

VISIT...

LANZAROTE
Caliente.com

К. Шеннон

РАБОТЫ ПО ТЕОРИИ ИНФОРМАЦИИ И КИБЕРНЕТИКЕ

ИЗДАТЕЛЬСТВО ИНОСТРАННОЙ ЛИТЕРАТУРЫ, Москва 1963

Книга представляет собой сборник статей выдающегося математика и инженера, члена Национальной академии наук США, Клода Эльвуда Шеннона. Многие из включенных в сборник работ, опубликованных проф. Шенноном в различных журналах в 1938—1962 годах, положили начало новым областям исследований в области общей теории связи, теории автоматов, электротехники, теории информации и лингвистики, таким, как теория анализа и синтеза релейных устройств, теория вероятностных схем, теория передачи информации и т. д.

Статьи расположены в сборнике по тематическому принципу: в первой части помещены работы по теории управляющих систем, во второй — по теории информации, в третьей — все остальные. В конце книги приводится библиография работ по теории информации.

Книга представляет интерес для широкого круга математиков и специалистов, работающих в области автоматического управления, теории связи, радиотехники, теории надежности и в смежных областях, так или иначе связанных с использованием результатов теории информации. Она будет полезна также студентам старших курсов университетов и технических вузов, инженерам и научным работникам различных специальностей, занимающимся вопросами, связанными с математическими аспектами кибернетики.

СОДЕРЖАНИЕ

Предисловие	5
ТЕОРИЯ УПРАВЛЯЮЩИХ СИСТЕМ	
Символический анализ релейных и переключательных схем	9
Число двухполюсных параллельно-последовательных сетей	46
Синтез двухполюсных переключательных схем	59
Требования, предъявляемые к объему памяти телефонного коммутатора	106
Надежные схемы из ненадежных реле	114
Использование машины для проектирования переключательных схем	154
Вычислительные устройства и автоматы	162
Машина для игры в шахматы	181
Составление программ для игры в шахматы на вычислительной машине	192
Играющие машины	216
Сообщение о машине, решающей лабиринтную задачу	223
Вклад фон Неймана в теорию автоматов	232
ТЕОРИЯ ИНФОРМАЦИИ	
Математическая теория связи	243
Теория связи в секретных системах	333
Современные достижения теории связи	403
Принципы кодово-импульсной модуляции	414
Связь при наличии шума	433
Некоторые задачи теории информации	461
Пропускная способность канала с шумом при нулевой ошибке	464

Геометрический подход к теории пропускной способности каналов связи	488
Каналы с дополнительной информацией на передатчике	497
Некоторые результаты теории кодирования для каналов с шумами	509
Замечания о частичном упорядочении каналов связи	532
Вероятность ошибки для оптимальных кодов в гауссовском канале	549
Теоремы кодирования для дискретного источника при заданном критерии точности	587
Двусторонние каналы связи	622
РАЗНОЕ	
Бандагон	667
Предсказание и энтропия печатного английского текста	669
Упрощенный вывод линейной теории сглаживания и предсказания по методу наименьших квадратов	687
Математическая теория дифференциального анализатора	709
О максимальном потоке через сеть	729
Теорема о раскраске ребер графа	735
Универсальная машина Тьюринга с двумя внутренними состояниями	740
Вычислимость на вероятностных машинах	751
Библиография	783
Именной указатель	821
Предметный указатель	824

ИМЕННОЙ УКАЗАТЕЛЬ

Аккерман (Ackermann W.) 13	Вернам (Vernam G. S.) 346
Александров А. Д. 488	Винер (Wiener N.) 180, 215, 293, 294, 309, 322, 403, 687
Александров П. С. 488	Виньерон (Vigneron H.) 215
Батлер (Butler S.) 162, 179	Виттстон (Whittstone) 347
Беннет (Bennet W.) 413, 437, 588	Вуджер (Woodger J.) 59
Беркли (Berkeley E. C.) 59, 179	Габор (Gabor D.) 403, 437
Беркс (Burks A.) 233	Галлагер (Gallager R. G.) 641
Беттел (Battel) 193	Гейне (Gaines H. F.) 333
Бигелоу (Bigelow) 226, 224, 231	Гельфонд А. О. 248
Биркгоф (Birkhoff G.) 59, 293	Герьери (Guerrieri J.) 711
Блекуэлл (Blackwell D.) 650, 685	Гилберт (Gilbert E. N.) 104, 574
Боде (Bode H. W.) 322, 463, 687	Гилиге (Gilliege M.) 333
Больцман (Boltzmann L.) 261, 403	Гильберт (Hilbert D.) 13, 56, 715
Брейман (Breiman L.) 658	Гобсон (Hobson E. W.) 782
Брозин (Brosin) 230	Голдстейн (Goldstine H.) 233
Буль (Bool G.) 13, 59	Гoley (Goley M. J. F.) 290
Бупш (Bush V.) 709	Гольден (Holden) 715
Бэббидж (Babbage Ch.) 162, 163	Греа (Grea R.) 102
Васильев Ю. Л. 102	Гупта (Gupta) 51, 53, 55
Ватанабе (Watanabe S.) 750	Гуревич (Hurewicz W.) 293, 445
Велч (Welch B. L.) 550, 551, 572	

- Давенпорт (Davenport W. B.) 291
 Данциг (Danzig G.) 730
 Девид (David H.) 558, 571
 Девис (Davis M.) 781
 Де Гроот (De Groot A. D.) 186, 196, 209, 215
 Де-Лей (De Leeuw K.) 571
 Дерр (Derr) 196
 Джерард (Gerard) 225, 231
 Джонсон (Johnson N.) 550, 551, 572
 Дибольд (Diebold J.) 179
 Добрушин Р. Л. 281, 319, 330, 543, 587, 598
 Дуб (Doob J. L.) 293, 343, 781
 Дьюи (Dewey G.) 253, 671
 Игонне (Higonnet R.) 102
 Кантор (Cantor G.) 445
 Кардо (Cardot C.) 102
 Кейстер (Keister W.) 219
 Кельвин (Kelvin) 709
 Кемпелен (Kempelen von W.) 182, 193, 217
 Кендалл (Kendall M. G.) 251
 Кёниг (Konig D.) 735
 Клини (Kleene S. C.) 165, 179, 782
 Колмогоров А. Н. 332, 343, 587, 687
 Кондон (Condon) 196
 Котельников В. А. 295, 435
 Котурат (Couturat L.) 13, 59
 Крамер (Cramer H.) 638
 Кричевский Р. Е. 58
 Крускал (Kruskal W.) 558, 571
 Крылов А. Н. 709
 Кудрявцев Л. Д. 735
 Купман (Koorman) 293
 Курант (Courant R.) 56
 Кэли (Cayley A.) 46
 Кэннон (Cannon W. B.) 175
 Ландау (Landahl H. D.) 179
 Литтлвуд (Littlewood J. E.) 682
 Лупанов О. Б. 102, 104
 Льюис (Luce R. D.) 341
 Любич Ю. И. 322
 Мак-Каллок (McCulloch W.) 148, 163, 179, 231
 Мак-Карти (McCarthy J.) 233
 Мак-Кинси (McKinsey J. C. C.) 341
 Мак-Коллум (McCullum D. M.) 167, 179
 Мак-Лейн (MacLane S.) 59
 Мак-Магон (Mac Mahon P.) 46, 50
 Мак-Миллан (McMillan B.) 322, 452
 Мецар (Meszar J.) 179
 Мид (Mead) 226
 Монтгомери (Montgomerie G.) 60
 Моргенштерн (Morgenstern O.) 194, 215, 341
 Мур (Moore E. F.) 114, 154, 169, 220
 Мурора (Muroga S.) 288, 488
 Мурский В. Л. 21
 Мюирхед (Muirhead) 682
 Найквист (Nyquist H.) 243, 403, 434, 437
 Накасима (Nakasima A.) 9, 60
 Нейман фон (Naumann von J.) 114, 176-178, 180, 194, 215, 231, 233—239, 293, 341
 Новиков П. С. 13
 Оливер (Oliver B.) 322, 414, 686
 Оруэлл (Orowell G.) 230
 Оттингер (Oettinger A. E.) 173, 179
 Петерсен (Petersen J.) 736
 Пиз (Pease W.) 180
 Пинскер М. С. 587
 Пирс (Pierce J. R.) 322, 414
 Питтс (Pitts W.) 59, 148, 179, 224
 Пиш (Piesch H.) 60
 Поваров Г. Н. 102, 103
 Пойа (Полна) (Polya G.) 682
 Пост (Post E. L.) 782
 Пратт (Pratt F.) 253, 366, 382, 670
 Раис (Rice S. O.) 542, 573
 Райт (Wright E. M.) 50, 195, 215
 Райффа (Raiffa H.) 341
 Риордан (Riordan J.) 46, 82, 83
 Россер (Rosser J.) 13
 Рут (Root W. L.) 291

Самуэль А. Л. (Samuel A. L.) 220
 Сарымсаков Г. А. 256
 Севидж (Savage L. J.) 223, 225, 226, 228 231
 Силверман (Silverman R.) 538 Смит (Smith J. B.) 167, 179, 251
 Стрэнч (Strachey C. S.) 170, 171, 180, 220, 309
 Сулливан (Sullivan H.) 309, 403
 Таллер (Tuller W. G.) 309, 403
 Тихомиров В. 332
 Тойбер (Teuber) 226, 230
 Толмен (Tolman R. C.) 261
 Томасян (Thomasian A.) 658
 Торрес-и-Квеведо (Torres y Quevedo L.) 182, 191, 218
 Тоуни (Tawney) 196
 Трахтенброт Б. А. 165
 Троттер (Trotter H.) 771
 Туркетт (Turquette A.) 13
 Тьюринг (Turing A. M.) 165, 166, 180, 236, 238, 740, 741, 782
 Уайтхед (Whitehead A.) 13
 Уивер (Weaver E.) 243
 Уиттекер (Whittaker J.) 437
 Уолман (Wallman H.) 445
 Фаддеев В. К. 261
 Фаин (Fine R.) 209, 215
 Файнштейн (Реш51еш А.) 281, 531, 729
 Фалкерсон (Fulkerson D.) 730
 Фано (Fano R. M.) 272
 Фелдман (Feldman C. B.) 413
 Феллер (Feller W.) 256, 257, 781
 Фёрстер (Foerster von) 228, 230
 Финк (Fink D. G.) 413

Форд (Ford L.) 730
 Фостер (Foster R. M.) 46, 736
 Франк (Frank) 226, 230
 Фреше (Frechet M.) 256
 Фробениус (Frobenius G.) 525, 531
 Халмош (Halmos P. R.) 781
 Хантингтон (Huntington E. V.) 14
 Харди (Hardy G. N.) 50, 195, 215, 682
 Харкевич А. А. 243
 Харкнес (Harkness) 193
 Хартли (Hartley R. V. L.) 243, 244, 403, 413, 434
 Хартри (Hartree D.) 709
 Хаусхолдер (Householder A. S.) 179
 Хегельбергер (Hagelbarger D. W.) 174, 221, 651, 652
 Хилл (Hill L. S.) 347
 Хинчин А. Я. 261, 270
 Холбрук (Holbrook B. D.) 110
 Хопф (Hopf E.) 293, 398,
 Хопф (Hopf H.) 488
 Хоэл (Hoel P. G.) 707
 Хэмминг (Hamming R.) 289
 Ципф (Zipf C. V.) 672
 Чандрасекар (Chandrasekhar S.) 250
 Чернев (Chernev) 215
 Чернов (Chernoff H.) 517, 531, 660
 Черч (Church A.) 781
 Шапиро (Shapiro N.) 751
 Шестаков В. И. 9
 Эйлер (Euler L.) 51
 Элайес (Elias P.) 487, 531, 543, 729
 Эшби (Ashby W. R.) 175, 176, 179
 Яблонский С. В. 13, 105
 Яглом А. М. 391
 Яглом И. М. 391

ПРЕДМЕТНЫЙ УКАЗАТЕЛЬ

Алгебра Буля 13, 59
 — — интерпретация 60
 — — класс функций 92, 128
 — — логическая интерпретация 13
 — — особые правила 61
 — — функция кворума 128
 — секретных систем 350

Анализатор контактных схем 154
 — — — особенности 161
 Ансамбль функций 291
 — — с ограниченной длительностью 295
 — — — — полосой частот 295
 — — «скорость числа измерений» 33

- — стационарный 292
- — эргодический 293
- Асимптотическая надежность 577
- Асимптотическое распределение 705
- Бит 244
- Вероятность ошибки 467, 514, 548
- — оптимальная 542, 551
- — границы 565, 566, 571, 574, 577
- — средняя для множества кодов 511
- — — ансамбля пар кодов 632
- Выходные последовательности случайного устройства 772
- Вычислительные устройства и автоматы 162
- Групповая инвариантность 93
- — для двух переменных 98
- — — трех переменных 98
- — специальный случай 94
- Декодирование 500
- по минимуму расстояния и максимуму правдоподобия 541
- Декодирующая система 465
- — оптимальная 541
- Дискретный преобразователь 268
- Дискретный преобразователь вырожденный 268
- Дифференциальный анализатор 709
- — автоматическое управление работой 726
- — идеализированный 711
- — передаточные числа 725
- Избыточность кода 267
- английского языка 267
- Инвариантное устройство (оператор) 294
- Информация взаимная 630
- мера 245
- скорость передачи 277, 305, 595
- средние взаимные скорости передачи 628
- Искажение 275
- полное 589

- Искажения мера 589, 612, 614. 616
- — схем 41
- Источник информации 245, 433
- — дискретный 249, 457
- — — с конечным числом состояний 264
- — непрерывный 458
- — произведение 596
- — скорость создания сообщений 318, 320, 458
- — смешанный 258
- — с независимыми буквами 591
- — эргодический 257, 325
- Исчисление высказываний 14
- Канал 246, 434
- гауссовский 540
- двоичный двусторонний умножающий 651
- — симметричный 543, 611
- — стирающий 543
- двусторонний 622, 627
- — с памятью 658
- — — симметричной структурой 647
- Канал дискретный без памяти 465, 509
- — — шума 247
- — — с шумом 275
- — непрерывный 305
- — произведение 466, 524
- с дополнительной информацией 497
- — конечным числом состояний 522, 527
- — обратной связью 479
- — памятью 620
- — Г концами 660
- сумма 564, 524
- частичное упорядочивание 532
- чистый 534
- Канала надежность 543
- Квантование 415
- Код блоковый для дискретного канала без памяти 509

- — — двустороннего канала 627
- — — канала с дополнительной информацией 499
- входная скорость 465
- райсовский 573
- случайный 470
- Хегельбергера 651
- Кодирование 604, 630, 631
- Кодовое слово 509 Крптограмма 335, 389
- Логика символическая 13
- — — система аксиом 14
- Матрица смежности 469
- Машина Тьюринга 165, 740
- — — метод построения 741
- — — моделирование 747
- — — понятие A -вычислимости 166
- — — проблема работы 742
- — — с двумя состояниями 741
- — — одним состоянием 744
- — — универсальная 166
- — — свойства 166
- Машины вероятностные 751
- — — классы 752
- — — вырабатывающие свои принципы игры 220
- — — вычислительной и мозга сравнение 237
- — — для игры в полностью проанализированные игры 218
- — — — полный анализ которой неизвестен 219
- — — нахождения пути в лабиринте 172
- Машины играющие 168, 216
- — — группы 168
- — — типы 218
- — — логические 162
- — — обучаемые 171
- — — гомеостат 175
- — — стохастические 765
- Модуляция амплитудная 420
- — — однополосная 442
- — — кодово-импульсная 403, 414, 425, 429
- — — фазово-импульсная 403, 452
- — — частотная 403, 425
- Ненадежность 365
- — — ключа 368, 370
- — — и избыточность 384
- — — распределение 385
- — — свойства 370
- — — случайного шифра 374
- — — сообщения 368, 371
- «Отклонение» вероятности 519
- Отношение сигнал/шум 423
- Отображение, понижающее смежность 474
- Переключатели селекторные 26
- — — шаговые 26
- Помехоустойчивость 424
- Пороговая мощность 419
- Пороговой эффект 445
- Предсказание английского языка 673
- — — идеальное n -граммное 677, 681
- — — линейное 689, 701, 708
- — — минимально-квадратичное 707
- — — «обратное» 677
- — — по методу Винера — Колмогорова 688
- — — при наличии шума 701
- Предсказания средне-квадратичная ошибка 694
- Пропускная способность, геометрический подход 488
- — — дискретного канала без памяти с дополнительной информацией 488, 501
- — — — — шума 247
- — — — — с шумом 281
- — — — — канала с конечным числом состояний 522, 527
- — — — — шумом при нулевой ошибке 464, 466, 467, 469
- — — — — памятью 621
- — — непрерывного канала 306

- Пропускная способность
 - непрерывного канала при ограничении средней мощности 308
 - — — — — пиковой мощности 309
 - — при наличии белого шума 446
 - — произвольном типе шума 455
- Пропускной способности область 628, 654, 659
- Процесс марковский 255, 342
- Рабочее число 33
- Реле идеализированные 116
 - нижняя оценка числа контактов 142
 - распределение нагрузок 66, 83
 - с неопределенным временем срабатывания 146
 - типы неисправностей 116
- Сглаживание 687, 703
 - линейное 689
- Секретные системы 333
 - — идеальные 338, 382
 - — идемпотентные 353
 - — подобные 337, 359
 - — совершенно 361, 363
 - — чистые 336
 - — эндоморфные 352
 - — «алгебраические системы» 353
- Система связи 245, 433
 - — геометрическое представление 437, 439, 441, 442
 - — дискретная 246
 - — идеальная 451
 - — непрерывная 246
 - — смешанная 246
 - — декодирования 509, 627
 - — рабочих чисел 33
- Сети аппроксимация числа 57
 - асимптотическое поведение 54
 - максимальный поток через 729
 - параллельно-последовательные 46
 - поведение 47
 - подсчет по Гупта 52
 - — — Дарбу 51
 - — — Эйлеру 52
 - получение производящей функции 48
 - правило соответствия 49
 - приведенные 731
 - расширенные 734
 - ребра 729
 - сечение 731
- Сети существенно-параллельные 47
 - существенно-последовательные 47
 - эквивалентность соединений 47
- Сопrotivления способы соединения 12
 - схем способы определения 65
 - функция 11, 21, 61
- Сумматор 44
- Схем общая теория 19
 - переключательных анализ 59
 - — синтез 59
 - — теория 59
 - сопротивление 10
- Схема гамакообразная 138, 144, 149
 - универсальная 38
- Схемы вероятность замыкания 122
 - — размыкания 121
 - верхние границы вероятности ошибок 149
 - двойное сопротивление 10
 - двойственные 30
 - двухполюсные 10, 729
 - длина 121
 - заданной длины и ширины 130
 - метод каскадов 100
 - — нахождения 31
 - — повышения надежности 117
 - — построения 21
 - мостиковые 46
 - направленного действия 148
 - нахождение двойственной 31
 - не параллельно-последовательные 21

- — — — методы построения 21
- — — — многополюсные 19
- параллельно-последовательные 17
- параллельные 11
- переключаемые 9, 154
- преобразования 19
- проектирование двухполюсных 155
 - на анализаторе 155
- разделительное дерево 89
- разделительные 67
- реализующей функцию синтез 32
- релейные 9
 - синтез 47
- с постоянным напряжением 19
- с блокировкой 26
- секторные 41
- символический анализ 9
- синтез надежных 147
 - переключательных 47
- символический анализ 9
- сложение сопротивлений 11
- с постоянным током 30
- Схемы типы 158
 - соединений 9
 - умножение сопротивлений 11
 - управления 9
 - свойства 9
 - формирование функций нескольких переменных 145
- функция $h(p)$ 120
- ширина 121
- эквивалентной данной, методы нахождения 9
- Теорема де Моргана 14
 - для канала без шума, основная 270
 - дискретного канала с шумом, основная 281
 - о раскраске ребер графа 735
 - схемах и функциях 26—30
 - способах соединения сопротивлений 12, 13
 - эргодических источниках 324
 - отсчетов (Котельникова) 435
- синтеза основная 67
- Теория автоматов 231
 - вклад фон Неймана 231
 - — — — «мультиплекс» 235
 - проблемы эволюции 237
- Точность передачи 315
- Точности критерий 316, 317
- Фильтр линейный 690
 - с минимальной фазовой характеристикой 693, 697
- Функции несимметрические 33
 - произвольной реализация 35
 - частично-симметрические, 100
 - $h(p)$ свойства 120
 - $\lambda(n)$ оценка 82
- Функций ансамбль 291
 - метод реализации 45
- Функций реализация 715
 - симметрических 35, 45
- Функциональная разделимость 93
- Функциональных соотношений типы 93
- Функция переключательная 58
 - порождающая 50
 - производящая 46, 50
 - получение 48
 - сопротивления схемы 61
 - $\lambda(n)$ 66
 - поведение 68
 - $\mu(n)$ 66
- Шифр Бофора 345, 360
- Вернама 346, 364
- Виженера 345, 346, 360, 558
- дробный 348
- Плэйфер 347
- простой подстановки 344
- случайный 374
- чистый 354, 356
- Цезаря 345, 358, 360, 366
- Шум белый 292, 304, 408, 409, 448, 454
 - мощность 456
 - Гауссовский, см. шум белый

— — произвольный 454
— квантования 415, 423
— тепловой, см. шум белый
Энтропийная мощность 301, 303, 456
Энтропия 261
— английского языка 670

— ансамбля функций 300, 303
— непрерывного распределения 296
— относительная 267
— условная 263, 508
Эффективный процесс 752

В наш век возрастающей дифференциации человеческих знаний Клод Шеннон является исключительным примером соединения глубины отвлеченной математической мысли с широким и в то же время совершенно конкретным пониманием больших проблем техники. Его в равной мере можно считать одним из первых математиков и одним из первых инженеров последних десятилетий. Своеобразная роль ему принадлежит в создании кибернетики. В отличие от Норберта Винера Шеннон не занимался пропагандой и систематизацией этой новой науки. Но он создал основы теории информации и в значительной мере предопределил своими работами развитие общей теории дискретных автоматов, которые составляют две большие главы кибернетики, занимающие в ней едва ли не центральное положение.

Значение работ Шеннона для чистой математики не сразу было достаточно оценено. Мне вспоминается, что еще на международном съезде математиков в Амстердаме (1954 г.) мои американские коллеги, специалисты по теории вероятностей, считали мой интерес к работам Шеннона несколько преувеличенным, так как это более техника, чем математика. Сейчас такие мнения вряд ли нуждаются в опровержении.

Правда, строгое математическое «обоснование» своих идей Шеннон в сколько-либо трудных случаях предоставил своим продолжателям. Однако его математическая интуиция изумительно точна. Мне известен только один случай, где она его, по-видимому, обманула: правильность формулы для λ в конце приложения 7 к работе «Математическая теория связи» вызывает сомнение.

Агитировать в среде специалистов по технике связи за внимание к работам Шеннона сейчас, можно думать, тоже излишне. Полное издание всех его работ на русском языке в высшей степени своевременен. Так как издание рассчитано на подготовленного читателя,

редакторы ограничились самыми необходимыми краткими примечаниями. С точки зрения общих перспектив развития теории информации и кибернетики может представлять интерес небольшая заметка самого Шеннона «Бандвагон». С скромный и деловой подход к имеющимся в этих областях к настоящему времени достижениям типичен для Шеннона.

Статьи, включенные в сборник, разделены на три раздела. Первый раздел и заключительные пять статей третьего раздела публикуются под редакцией О. Б. Лупанова; второй раздел, три первые статьи третьего раздела и библиография — под редакцией Р. Л. Добрушина.

А. Колмогоров

Теория управляющих систем

1. Введение

В схемах управления и защиты сложных электрических систем часто бывают необходимы сложные соединения контактов реле и переключателей. Такие схемы используются в автоматических телефонных коммутаторах, аппаратуре управления двигателями и в большинстве схем, предназначенных для автоматизации сложных процессов. В этой статье излагается математический анализ некоторых свойств таких схем. Особое внимание уделяется проблеме синтеза схем.

Пусть даны некоторые условия функционирования; требуется найти схему, реализующую эти условия. Решение задач такого рода неоднозначно; в работе исследуются методы нахождения индивидуальных схем, требующих наименьшего числа контактов реле и переключателей. Также дается описание методов нахождения схем, эквивалентных данной по всем заданным условиям функционирования. В работе показано, что некоторым хорошо известным теоремам о схемах, составленных из элементов с заданным импедансом, соответствуют аналогичные теоремы о релейных схемах. Отметим здесь теоремы о преобразованиях треугольника в звезду и звезды в ячейку и теорему двойственности.

Метод подхода к решению поставленных проблем может быть кратко описан следующим образом. Любая схема представляется в виде системы уравнений, составленных из символов, соответствующих различным реле и переключателям схемы. Разрабаты-

¹⁾ Shannon C., A Symbolic Analysis of Relay and Switching Circuits, *Transactions of the American Institute of Electrical Engineers*, 57 (1938), 713—723. [Рукопись сдана 1 марта 1938 г., подготовлена к печати 27 мая 1938 г. Статья является рефератом диссертации, представленной в Массачусетский технологический институт на соискание степени магистра.]

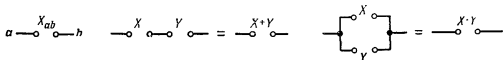
²⁾ Почти одновременно с данной работой советским ученым В. И. Шестаковым (1935 г.; опубликовано в 1941 г.) и японским ученым А. Накаси́ма (1938 г.) было установлено, что контактные схемы описываются с помощью функций алгебры логики. (См. Шестаков В. И., Алгебра двухполюсных схем, построенных исключительно из двухполюсников (алгебра А-схем), *Ж. Техн. Физ.*, 11:6 (1941); Nakasima A., Цикл статей в журнале *Nippon electr. Comm. of Japan* за 1938 г.) — *Прим. ред.*

вается аппарат для преобразования этих уравнений с помощью простых математических приемов, большинство из которых подобно обычным алгебраическим алгоритмам. Показывается, что этот аппарат в точности аналогичен исчислению высказываний символической логики. Для синтеза схемы заданные условия сначала записываются в виде системы уравнений, затем уравнения преобразуются к виду, соответствующему простейшей схеме. Тогда схема может быть получена непосредственно из уравнений. Этим методом всегда можно найти простейшую параллельно-последовательную схему, а в некоторых случаях — простейшую схему, содержащую любые типы соединений.

Используемые нами обозначения взяты главным образом из символической логики. Из большого многообразия применяемых в наши дни систем выбрана та, которая представляется более простой и удобной для нашей интерпретации. Некоторые из употребляемых нами терминов, как, например, узел, ячейка, треугольник, звезда и т. п., заимствованы из общей теории электротехнических схем для обозначения сходных понятий в переключательных схемах.

2. Параллельно-последовательные двухполюсные схемы. Основные определения и постулаты

Ограничим наше исследование схемами, содержащими только контакты реле и переключателей. Такие схемы между любыми двумя своими полюсами в каждый момент времени либо замкнуты (нулевое сопротивление), либо разомкнуты (бесконечное сопротивление). Сопоставим полюсам a и b символ X_{ab} или проще X . Эту



Р и с. 1. Символическое изображение функции сопротивления.

Р и с. 2. Интерпретация сложения.

Р и с. 3. Интерпретация умножения.

переменную функцию времени будем называть двоичным сопротивлением двухполюсной схемы $a-b$. Символом 0 (нуль) будем обозначать сопротивление замкнутой схемы, а символом 1 (единица) — сопротивление разомкнутой схемы. Следовательно, если схема $a-b$ разомкнута, то $X_{ab} = 1$, а если замкнута, то $X_{ab} = 0$. Сопротивления X_{ab} и X_{cd} называются равными, если схема $a-b$ разомкнута тогда и только тогда, когда схема $c-d$ разомкнута. Пусть теперь символ $+$ (плюс) определяется в смысле последовательного соединения двухполюсных схем, при котором сопротивления скла-

дываются. Так $X_{ab} + X_{cd}$ — это сопротивление схемы $a - d$, полученной в результате объединения полюсов b и c схем $a - b$ и $c - d$. Аналогично произведение двух сопротивлений $X_{ab} \cdot X_{cd}$, или короче $X_{ab}X_{cd}$, определяется как сопротивление схемы, образованной параллельным соединением схем $a - b$ и $c - d$ ¹⁾. Контакты реле или переключателей изображаются в схемах, как показано на рис. 1, где буквы — это соответствующие функции сопротивления. На рис. 2 показана интерпретация сложения, а на рис. 3 — умножения. Такой выбор символики делает операции над сопротивлениями очень похожими на обычные алгебраические преобразования.

Очевидно, что при приведенных выше определениях справедливы следующие постулаты.

1. $a. 0 \cdot 0 = 0$ Замкнутая схема, соединенная параллельно с замкнутой схемой, есть замкнутая схема.
 $b. 1 + 1 = 1$ Разомкнутая схема, соединенная последовательно с разомкнутой схемой, есть разомкнутая схема.
2. $a. 1 + 0 = 0 + 1 = 1$ Разомкнутая схема, соединенная последовательно с замкнутой схемой в любом порядке (т. е. разомкнутая схема справа или слева от замкнутой), есть разомкнутая схема.
 $b. 0 \cdot 1 = 1 \cdot 0 = 0$ Замкнутая схема, соединенная параллельно с разомкнутой схемой (в любом порядке), есть замкнутая схема.
3. $a. 0 + 0 = 0$ Замкнутая схема, соединенная последовательно с замкнутой схемой, есть замкнутая схема.
 $b. 1 \cdot 1 = 1$ Разомкнутая схема, соединенная параллельно с разомкнутой схемой, есть разомкнутая схема.
4. В каждый момент либо $X = 0$, либо $X = 1$.

Этих постулатов достаточно, чтобы вывести все теоремы, которые будут использованы в связи со схемами, образованными только из последовательных и параллельных соединений. Постулаты расположены попарно, чтобы подчеркнуть двойственность между операциями сложения и умножения и константами 0 и 1. Так, если

¹⁾ В нашей литературе функционирование контактных схем чаще описывается двойственным образом — с помощью функций проводимости: замкнутая схема имеет проводимость 1, разомкнутая — 0; при параллельном соединении схем их проводимости (логически) складываются, а при последовательном — умножаются. — *Прим. ред.*

в каком-нибудь постулате, относящемся к типу *a*, нули заменить на единицы, умножение на сложение и наоборот, то получится соответствующий постулат типа *b*. Этот факт очень важен. Он дает для каждой теоремы двойственную; достаточно доказать одну, чтобы установить обе. Единственный из введенных нами постулатов, отличающийся от постулатов обычной алгебры, это *1b*. Тем не менее он дает возможность существенно упростить операции над нашими символами.

Теоремы

Приведем ряд теорем о способах соединения сопротивлений. Так как любая из этих теорем может быть доказана очень простым способом, приведем один пример доказательства в качестве иллюстрации. Методом доказательства является метод «полной индукции», т. е. метод проверки теоремы для всех возможных случаев. Так как в силу постулата 4 каждое переменное может принимать только значения 0 и 1, это легко сделать. Некоторые теоремы могут быть доказаны более изящно сведением к предыдущим теоремам, но метод полной индукции настолько универсален, что, вероятно, ему следует отдать предпочтение.

$$X + Y = Y + X; \quad (1a)$$

$$XY = YX; \quad (1b)$$

$$X + (Y + Z) = (X + Y) + Z; \quad (2a)$$

$$X(YZ) = (XY)Z; \quad (2b)$$

$$X(Y + Z) = XY + XZ; \quad (3a)$$

$$X + YZ = (X + Y)(X + Z); \quad (3b)$$

$$1 \cdot X = X; \quad (4a)$$

$$0 + X = X; \quad (4b)$$

$$1 + X = 1; \quad (5a)$$

$$0 \cdot X = 0. \quad (5b)$$

Например, чтобы доказать теорему 4а, заметим, что *X* есть либо 0, либо 1. Если *X* = 0, теорема следует из постулата 2b; если *X* = 1, она следует из постулата 3b. Теорема 4b теперь вытекает из 4а по принципу двойственности в результате замены 0 на 1 и $\cdot$ на $+$.

В силу ассоциативных законов (2а) и (2b) в сумме или произведении нескольких членов скобки могут быть опущены. Символы Σ и Π имеют те же значения, что и в обычной алгебре.

Дистрибутивный закон (3а) дает возможность «развернуть» умножение и «свернуть» сумму. Двойственная теорема (3b), однако, в обычной алгебре неверна.

Определим теперь новую операцию, которую назовем отрицанием. Отрицание сопротивления X обозначается через X' и определяется как переменная, равная 1, когда X равно 0, и равная 0, когда X равно 1. Если X — сопротивление замыкающего контакта реле, то X' — сопротивление размыкающего контакта того же реле. Из определения отрицания сопротивления вытекают теоремы:

$$X + X' = 1; \quad (6a)$$

$$XX' = 0; \quad (6b)$$

$$0' = 1; \quad (7a)$$

$$1' = 0; \quad (7b)$$

$$(X')' = X. \quad (8)$$

Аналогия с исчислением высказываний

Теперь можно доказать эквивалентность введенного исчисления некоторой элементарной части исчисления высказываний. Алгебра логики¹⁾, введенная Джорджем Булем, является символическим методом вывода логических соотношений. Символы булевой алгебры допускают две логические интерпретации. При интерпретации в терминах классов переменные могут принимать не только два значения 0 и 1. Эта интерпретация называется алгеброй классов. Если, однако, символы рассматриваются как высказывания, то имеем исчисление высказываний, в котором переменные принимают только значения 0 и 1²⁾, как и рассмотренные выше функции сопротивления. Обычно обе интерпретации основываются на одном и том же множестве постулатов, за исключением того, что в случае

¹⁾ Полный список литературы по символической логике приведен в журнале *Journal of Symbolic Logic*, 1, 4, декабрь 1936. Части теории, связанные с релейными схемами, достаточно полно освещаются в работах: Couturat L., *The algebra of logic*, Open Court, 1914; Whitehead A. N., *Universal algebra*, Cambridge Univ. Press, v. I, b. III, ch. I, II. [Из доступной литературы по этому вопросу на русском языке следует указать гл. книги Д. Гильберта и В. Аккермана «Основы теоретической логики», ИЛ, 1947, гл. книги П. С. Новикова «Элементы математической логики», Физматгиз, 1959, и в особенности гл. статьи С. В. Яблонского «Функциональные построения в k -значной логике», *Труды Матем. ин-та им. В. А. Стеклова АН СССР*, 51 (1958). Последняя является хорошим пособием для первоначального ознакомления с основами алгебры логики и теории контактных схем. — Прим. ред.]

²⁾ Это относится только к классической теории исчисления высказываний. Недавно появился ряд работ по логическим системам, в которых высказывание может принимать более двух «значений истинности». [Из более поздних фундаментальных работ можно указать: Rosser J., *Turquette A.*, *Many-valued logics*, Amsterdam, 1952; Яблонский С. В., *Функциональные построения в k -значной логике*, *Труды Матем. ин-та им. В. А. Стеклова АН СССР*, 51 (1958), 5 — 142. — Прим. ред.]

исчисления высказываний к постулатам 1—4 добавляется постулат эквивалентности. Е. В. Хантингтон¹⁾ дает следующую систему аксиом символической логики:

1. Класс K содержит по крайней мере два различных элемента;
2. Если a и b принадлежат классу K , то и $a + b$ принадлежит классу K ;

$$3. a + b = b + a;$$

$$4. (a + b) + c = a + (b + c);$$

$$5. a + a = a;$$

$$6. ab + ab' = a, \text{ где } ab \text{ определяется как } (a' + b')'.$$

Если положить, что класс K состоит из двух элементов 0 и 1, то эти постулаты будут следствиями постулатов, приведенных во введении, и наоборот, приведенные там постулаты 1, 2, 3 могут быть выведены из постулатов Хантингтона. Если добавить постулат 4 и ограничиться исчислением высказываний, то становится очевидной полная аналогия между этой ветвью символической логики²⁾ и исчислением переключательных схем. Обе интерпретации символов показаны в табл. I.

Благодаря этой аналогии любая теорема исчисления высказываний является также истинной теоремой, если ее интерпретировать в терминах релейных схем. Остальные теоремы данного раздела были установлены непосредственно на этой основе.

Теоремы де Моргана

$$(X + Y + Z + \dots)' = X' \cdot Y' \cdot Z' \cdot \dots; \quad (9a)$$

$$(X \cdot Y \cdot Z \dots)' = X' + Y' + Z' + \dots \quad (9b)$$

выражают отрицание суммы или произведения в терминах отрицания слагаемых или сомножителей. Они могут быть проверены для двух членов подстановкой всех возможных значений, а затем методом индукции распространены на любое число n переменных.

Функция нескольких переменных $X_1, X_2, \dots, X_n$ — это выражение, образованное из переменных при помощи операций сложения, умножения и отрицания³⁾. Такая функция будет обозначаться через $f(X_1, X_2, \dots, X_n)$. Так, например, мы можем записать:

$$f(X, Y, Z) = XY + X'(Y' + Z').$$

¹⁾ Huntington E. V., *Trans. Amer. Math. Soc.*, 36 (1933), 274—304.

²⁾ Эту аналогию можно усмотреть с несколько иной точки зрения. Вместо того чтобы сопоставлять X_{ab} непосредственно схеме $a-b$, будем интерпретировать X_{1b} как высказывание «схема $a-b$ разомкнута». Тогда все символы непосредственно интерпретируются как высказывания, а операции сложения и умножения рассматриваются как последовательное и параллельное соединения.

³⁾ Автор не различает логическое выражение (формулу) и определяемую им функцию (соответствие между наборами значений переменных и значениями функции). — *Прим. ред.*

Таблица I

Аналогия между исчислением высказываний и символическим анализом релейных схем

Символы	Интерпретация в терминах релейных схем	Интерпретация в терминах исчисления высказываний
X	Схема X	Высказывание X
0	Схема замкнута	Высказывание ложно
1	Схема разомкнута	Высказывание истинно
$X + Y$	Последовательное соединение схем X и Y	Высказывание истинно, если либо X , либо Y истинно
XY	Параллельное соединение схем X и Y	Высказывание истинно, если как X , так и Y истинно
X'	Схема, которая разомкнута, если X замкнута, и замкнута, если X разомкнута	Отрицание высказывания X
$=$	Схемы разомкнуты или замкнуты одновременно	Каждое из высказываний влечет другое

В анализе бесконечно малых показано, что любую функцию (при условии, что она непрерывна и все ее производные непрерывны) можно разложить в ряд Тейлора. Некоторое подобное разложение возможно и в исчислении высказываний. Чтобы вывести разложение функции в ряд, рассмотрим сначала следующие равенства:

$$f(X_1, X_2, \dots, X_n) = X_1 f(1, X_2, \dots, X_n) + X'_1 f(0, X_2, \dots, X_n); \quad (10a)$$

$$f(X_1, X_2, \dots, X_n) = [f(0, X_2, \dots, X_n) + X_1] [f(1, X_2, \dots, X_n) + X'_1]. \quad (10b)$$

Они превращаются в тождества, если положить $X_1 = 0$ или $X_1 = 1$. В этих равенствах функция f называется разложенной по X_1 . Коэффициенты при X_1 и X'_1 в (10a) являются функциями от $n-1$

переменных и могут быть таким же образом разложены по любой из этих переменных. Аддитивные члены в (10b) также могут быть разложены таким же образом. Разлагая по X_2 , получим:

$$\begin{aligned} f(X_1, X_2, \dots, X_n) = & X_1 X_2 f(1, 1, X_3, \dots, X_n) + \\ & + X_1 X_2' f(1, 0, X_3, \dots, X_n) + \\ & + X_1' X_2 f(0, 1, X_3, \dots, X_n) + \\ & + X_1' X_2' f(0, 0, X_3, \dots, X_n); \end{aligned} \quad (11a)$$

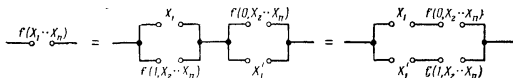
$$\begin{aligned} f(X_1, X_2, \dots, X_n) = & [X_1 + X_2 + f(0, 0, X_3, \dots, X_n)] \cdot \\ & \cdot [X_1 + X_2' + f(0, 1, X_3, \dots, X_n)] \cdot \\ & \cdot [X_1' + X_2 + f(1, 0, X_3, \dots, X_n)] \cdot \\ & \cdot [X_1' + X_2' + f(1, 1, X_3, \dots, X_n)]. \end{aligned} \quad (11b)$$

Повторяя разложение n раз, придем к полным разложениям в ряд, имеющим вид

$$\begin{aligned} f(X_1, X_2, \dots, X_n) = & f(1, 1, \dots, 1) X_1 X_2 \dots X_n + \\ & + f(0, 1, \dots, 1) X_1' X_2 \dots X_n + \dots + f(0, 0, \dots, 0) X_1' X_2' \dots X_n'; \end{aligned} \quad (12a)$$

$$\begin{aligned} f(X_1, X_2, \dots, X_n) = & [X_1 + X_2 + \dots + X_n + f(0, 0, \dots, 0)] \cdot \\ & \dots [X_1' + X_2' + \dots + X_n' + f(1, 1, \dots, 1)]. \end{aligned} \quad (12b)$$

Согласно (12a), f равна сумме произведений, полученных путем расстановки знаков отрицаний при $X_1, X_2, \dots, X_n$ всеми возможными способами и умножения каждого произведения на коэффициент, равный значению функции, когда это произведение есть 1. Аналогично для (12b).



Р и с. 4. Разложение по одной переменной.

В качестве приложения такого разложения покажем, что, если требуется найти схему, реализующую данную функцию, можно всегда разложить эту функцию по формуле (10a) или (10b) так, что некоторая выделенная переменная встречается не более двух раз — один раз как замыкающий и один раз как размыкающий контакт. Это показано на рис. 4. Согласно формулам (11a) и (11b) другая переменная встречается не более четырех раз (два раза как замыкающий контакт и два раза — как размыкающий) и т. д.

Обобщение теоремы де Моргана представляется символически следующим уравнением:

$$(f(X_1, X_2, \dots, X_n, +, \cdot))' = f(X'_1, X'_2, \dots, X'_n, \cdot, +). \quad (13)$$

Под этим подразумевается, что отрицание любой функции может быть получено заменой каждой переменной ее отрицанием и перестановкой символов $+$ и $\cdot$. Явные и неявные скобки остаются, конечно, на тех же местах. Например, отрицание $X + Y \cdot (Z + W \cdot X')$ будет иметь вид $X' [Y' + Z' (W' + X)]$.

Приведем некоторые другие теоремы, используемые для упрощения формул:

$$X = X + X = X + X + X = \dots; \quad (14a)$$

$$X = X \cdot X = X \cdot X \cdot X = \dots; \quad (14b)$$

$$X + XY = X; \quad (15a)$$

$$X(X + Y) = X; \quad (15b)$$

$$XY + X'Z = XY + X'Z + YZ; \quad (16a)$$

$$(X + Y)(X' + Z) = (X + Y)(X' + Z)(Y + Z); \quad (16b)$$

$$Xf(X, Y, Z, \dots) = Xf(1, Y, Z, \dots); \quad (17a)$$

$$X + f(X, Y, Z, \dots) = X + f(0, Y, Z, \dots); \quad (17b)$$

$$X'f(X, Y, Z, \dots) = X'f(0, Y, Z, \dots); \quad (18a)$$

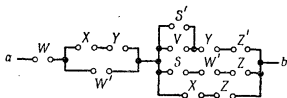
$$X' + f(X, Y, Z, \dots) = X' + f(1, Y, Z, \dots). \quad (18b)$$

Все эти теоремы могут быть доказаны методом полной индукции.

Любое выражение, образованное при помощи операций сложения, умножения и отрицания, является точным представлением схемы, содержащей только последовательные и параллельные соединения. Такая схема называется параллельно-последовательной. Каждая буква в выражении такого рода представляет замыкающий, размыкающий или переключающий контакт реле или переключателя. Поэтому, чтобы найти параллельно-последовательную схему, содержащую наименьшее число контактов, необходимо преобразовать выражение к форме, содержащей наименьшее число букв. Для этого вполне достаточно теорем, приведенных выше. Небольшой навык в оперировании символами — вот все, что требуется. К счастью, большинство из этих теорем в точности такие же, как в обычной алгебре. Автор считает, что особенно полезными для упрощения сложных выражений являются теоремы 3, 6, 9, 14, 15, 16а, 17 и 18.

Часто функция может быть записана различными способами, требующими одного и того же минимального числа элементов.

В таком случае может быть выбрана любая из схем или выбор может быть продиктован теми или иными соображениями.



Р и с. 5. Схема, подлежащая упрощению.

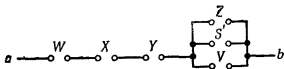
В качестве примера упрощения формулы рассмотрим схему, изображенную на рис. 5. Функцией сопротивления X_{ab} этой схемы будет:

$$\begin{aligned} X_{ab} &= W + W' (X + Y) + (X + Z) (S + W' + Z) (Z' + Y + S'V) = \\ &= W + X + Y + (X + Z) (S + 1 + Z) (Z' + Y + S'V) = \\ &= W + X + Y + Z (Z' + S'V). \end{aligned}$$

Эти преобразования были произведены с помощью формулы (17b), где в качестве X последовательно бралось сначала W , затем X и Y . Раскрывая скобки, получим

$$X_{ab} = W + X + Y + ZZ' + ZS'V = W + X + Y + ZS'V.$$

Схема, соответствующая этой формуле, изображена на рис. 6. Следует обратить внимание на большое сокращение числа элементов.



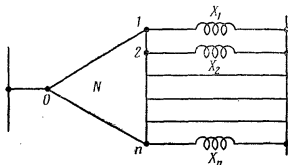
Р и с. 6. Упрощение схемы, изображенной на рис. 5.

При графическом представлении схем обмотки реле удобно обозначать той же буквой, что и сопротивление замыкающего контакта этого реле. Так, если обмотка реле соединена с источником питания схемой, функция сопротивления которой есть X , то само реле и любой его замыкающий контакт обозначается через X . Размыкающий контакт обозначается через X' . При этом предполагается, что реле срабатывает мгновенно и что одновременно замыкающий контакт замыкается и размыкающий размыкается. Случай, когда имеется задержка во времени, будет рассмотрен ниже.

3. Многополюсные и не параллельно-последовательные схемы

Эквивалентность n -полюсных схем

Обычная релейная управляющая схема имеет вид, изображенный на рис. 7, где $X_1, X_2, \dots, X_n$ — реле или другие устройства, управляемые схемой, а N — схема из контактов реле и переключателей. Требуется найти преобразования, применение которых к N



Р и с. 7. Общий вид релейной схемы с постоянным напряжением.

сохранит функционирование *всех* реле $X_1, X_2, \dots, X_n$. До сих пор рассматривались только преобразования, применимые к двухполюсным схемам и сохраняющие функционирование одного реле, последовательно соединенного со схемой.

Определим эквивалентность n -полюсных схем следующим образом.

О п р е д е л е н и е. Две n -полюсные схемы M и N называются эквивалентными относительно этих n полюсов, если $X_{jk} = Y_{jk}$, $j, k = 1, 2, \dots, n$, где X_{jk} — сопротивление схемы N (рассматриваемой как двухполюсная схема) между полюсами j и k , а Y_{jk} — сопротивление схемы M между соответствующими полюсами.

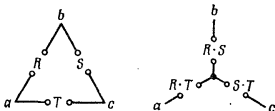
Под это определение подпадает и рассмотренное выше понятие эквивалентности для двухполюсных схем.

Преобразование треугольника в звезду и звезды в ячейку

Как и в общей теории схем, здесь существуют преобразования треугольника в звезду и звезды в ячейку. В схемах с конечным сопротивлением такие преобразования, если они существуют, единственны. В схемах рассматриваемого здесь типа эти преобразования всегда существуют, но не однозначны. Приводимые здесь схемы являются самыми простыми, так как они требуют наименьшего числа элементов. Преобразование треугольника в звезду показано на

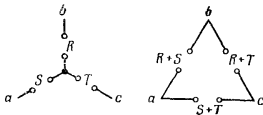
рис. 8. Обе схемы эквивалентны относительно трех полюсов a , b и c , так как по дистрибутивному закону $X_{ab} = R(S + T) = RS + RT$, и аналогичные соотношения имеют место для других пар полюсов $a-c$ и $b-c$.

Преобразование звезды в ячейку показано на рис. 9. Эквивалентность схем, изображенных на этом рисунке, следует из того



Р и с. 8. Преобразование треугольника в звезду.

факта, что $X_{ab} = R + S = (R + S)(R + T + T + S)$ и т. д. Для звезды с n лучами также существует эквивалентная ячейка. Эта ячейка получается удалением центра звезды и соединением



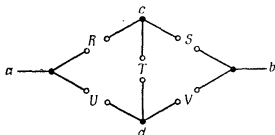
Р и с. 9. Преобразование звезды в ячейку.

каждой пары ее полюсов сопротивлением, равным сумме сопротивлений соответствующих лучей звезды. Это может быть доказано методом математической индукции. Было показано, что это верно для $n_1 = 3$. Предполагая это утверждение верным для $n-1$, докажем его для n . Предположим, что указанным методом построена ячейка из данной звезды с n лучами. Каждый угол полученной ячейки является звездой с $n-1$ лучом, и так как по предположению теорема верна для $n-1$, можно заменить n -й угол эквивалентной ему ячейкой. Если Y_{0j} — сопротивление первоначальной звезды от центра O до точки j , то полученная ячейка будет иметь между вершинами r и s сопротивление $(Y_{0s} + Y_{0r})(Y_{0s} + Y_{0n} + Y_{0r} + Y_{0n})$. Но это сводится к $Y_{0s} + Y_{0r}$, что является правильным значением, так как исходная звезда с n лучами, у которой удален n -й луч, является звездой с $n-1$ лучом и, по нашему предположению, может быть заменена ячейкой с этим сопротивлением между полюсами r и s . Следовательно, обе схемы эквивалентны относи-

тельно первых $n-1$ полюсов. При помощи элиминации другого полюса или из соображений симметрии доказывается эквивалентность относительно всех n полюсов¹⁾.

Функция сопротивления не параллельно-последовательной схемы

Методы, изложенные в разд. 2, недостаточны для описания и построения схем, содержащих не только параллельно-последовательные соединения. Примером схемы, не являющейся параллельно-последовательной, служит мостик, изображенный на рис. 10.



Р и с. 10. Схема, не являющаяся параллельно-последовательной.

Будем описывать такие схемы при помощи сведения их к эквивалентным параллельно-последовательным схемам. Разработано три метода построения схем, эквивалентных мостиковым схемам.

Первый — это очевидный метод применения преобразований к схеме до тех пор, пока не получится параллельно-последовательная схема с последующим вычислением соответствующей функции сопротивления. Это в точности тот же процесс, которым пользуются при упрощении сложных схем с конечным сопротивлением. Чтобы применить его к схеме, изображенной на рис. 10, сначала можно элиминировать узел c , применяя преобразование звезды $a-c$, $b-c$, $d-c$ в ячейку. Получим схему, изображенную на рис. 11. Функция сопротивления может быть найдена путем исследования полученной схемы

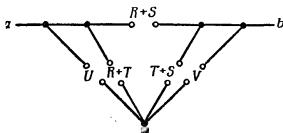
$$X_{ab} = (R + S) [U(R + T) + V(T + S)].$$

Это выражение можно преобразовать дальше так:

$$X_{ab} = RU + SV + RTV + STU = R(U + TV) + S(V + TU).$$

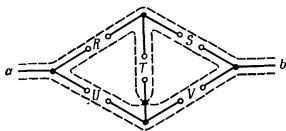
¹⁾ В. Л. Мурским построена полная система эквивалентных преобразований рассматриваемого автором вида для схем с любым числом полюсов, содержащих только контакты и не содержащих обмоток реле (Мурский В. Л., Об эквивалентных преобразованиях контактных схем, Сб. «Проблемы кибернетики», вып. 5, Физматгиз, 1961, 61—76). — Прим. ред.

Второй метод состоит в том, что в схеме выделяются все цепи, существующие между ее полюсами. Если сопротивление какой-нибудь из этих цепей равно нулю, искомая функция должна быть



Р и с. 11. Функция сопротивления, полученная в результате преобразований.

равна нулю. Следовательно, если результат записан в виде произведения, то сопротивление каждой цепи должно быть множителем этого произведения. Итак, требуемый результат может быть записан как произведение сопротивлений всех возможных цепей между



Р и с. 12. Функция сопротивления, представленная в виде произведения сумм.

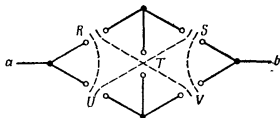
двумя полюсами. Цепи, которые проходят через один и тот же узел более одного раза, но нужно рассматривать. На рис. 12 показано применение этого метода к мостику. Цепи показаны пунктиром. Функция, следовательно, определяется формулой

$$\begin{aligned} X_{ab} &= (R + S) (U + V) (R + T + V) (U + T + S) = \\ &= RU + SV + RTV + UTS = \\ &= R (U + TV) + S (V + TU). \end{aligned}$$

Этот же результат был получен первым методом.

Третий метод состоит в том, что через контакты проводятся все возможные линии, пересекающие схему между рассматриваемыми полюсами. Результат записывается как сумма, каждый член которой соответствует некоторой из этих линий. Каждый такой член есть произведение сопротивлений всех контактов, расположенных на

этой линии. Обоснование этого метода подобно обоснованию второго метода. Применение его к мосту показано на рис. 13.



Р и с. 13. Функция сопротивления, представленная в виде суммы произведений.

При этом снова получаем

$$X_{ab} = RU + SV + RTV + STU = R(U + TV) + S(V + TU).$$

Третий метод является наиболее удобным и быстрым, так как он дает результат непосредственно в виде суммы. Представляется, несомненно, более легким иметь дело с суммами, чем с произведениями, так как в обычной алгебре имеем дистрибутивный закон $X(Y + Z) = XY + XZ$, но не имеем двойственного ему $X + YZ = (X + Y)(X + Z)$. Однако иногда затруднительно применять третий метод к неплоским схемам (т. е. к схемам, которые не могут быть изображены на плоскости без пересечений), и в этом случае может быть применен один из предыдущих методов.

Системы уравнений

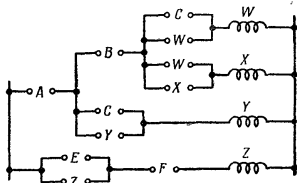
При анализе данной схемы удобно разбить различные переменные на два класса. Элементы, сопротивления которых управляются непосредственно источником, расположенным вне рассматриваемой схемы, будем называть независимыми переменными. К ним относятся ручные переключатели, контакты внешних реле и т. п. Реле и другие устройства, управляемые схемой, будем называть зависимыми переменными. Будем как правило, пользоваться первыми буквами алфавита для обозначения независимых переменных, а последними — для обозначения зависимых переменных. На рис. 7 зависимыми переменными являются $X_1, X_2, \dots, X_n$. Очевидно, реле X_k будет возбуждено тогда и только тогда, когда $X_{0k} = 0$, где X_{0k} — функция сопротивления схемы N между полюсами 0 и k . Это значит, что

$$X_k = X_{0k}, \quad k = 1, 2, \dots, n.$$

Эта система уравнений полностью определяет функционирование системы. Члены, стоящие в правых частях, должны быть извест-

ными функциями, включающими различные зависимые и независимые переменные. Значения зависимых переменных могут быть вычислены, если даны начальные условия и значения независимых переменных.

Опишем преобразования, уменьшающие число элементов, требующихся для реализации системы уравнений. Эти преобразования не меняют X_{0k} , $k = 1, 2, \dots, n$, но X_{jk} , $j, k = 1, 2, \dots, n$ могут меняться. Поэтому новая схема может не быть в строгом смысле



Р и с. 14. Пример упрощения системы уравнений.

эквивалентной старой, однако действие всех реле останется тем же. Этот метод упрощения применим только в том случае, когда функции X_{0k} записаны в виде сумм и некоторые члены являются общими для двух или большего числа уравнений. Например, предположим, что дана следующая система уравнений:

$$W = A + B + CW;$$

$$X = A + B + WX;$$

$$Y = A + CY;$$

$$Z = EZ + F.$$

Она может быть реализована схемой, изображенной на рис. 14, использующей только один элемент A при трех его вхождениях в систему и только один элемент B при двух его вхождениях. Обоснование почти очевидно, что может быть представлено графически проведением вертикальной черты после члена, общего для разных уравнений

$$\begin{aligned} W &= \\ X &= A + \left| B + \right| CW; \\ Y &= \left| CY; \right. \\ Z &= F + EZ. \end{aligned}$$

Из принципа двойственности следует, что, если последовательно-му соединению поставить в соответствие умножение, а параллельно-му — сложение, получаются в точности те же теоремы для преобразования. Имеются два соображения в пользу выбора приведенных определений. Во-первых, как было отмечено, легче оперировать с суммами, чем с произведениями, а только что описанное преобразование может быть применено только к суммам; во-вторых, при таком выборе функция сопротивления (в нашем смысле) в точности аналогична импедансу. При противоположном определении она была бы более близка к функции проводимости цепей переменного тока, которая используется реже.

Иногда между реле X и Y имеется соотношение $XY' = 0$. Оно имеет место, если реле Y может сработать, только если сработало реле X . Это часто имеет место в системах последовательного действия. В схемах такого типа реле могут срабатывать только в определенном порядке или последовательности, действие одного реле подготавливает схему к тому, чтобы могло сработать следующее по порядку реле. Если срабатывание реле X предшествует срабатыванию реле Y и далее оба реле остаются возбужденными, пока не работает вся последовательность, то эти условия выполнены. В этом случае для упрощений формул иногда могут быть использованы следующие равенства. Именно, если $XY' = 0$, то

$$X'Y' = Y';$$

$$XY = X;$$

$$X' + Y = 1;$$

$$X' + Y' = X';$$

$$X + Y = Y.$$

Нетрудно убедиться в справедливости этих равенств, умножив левую часть любого из них на $X' + Y$, равное 1, или добавив к ней XY' , равное 0. Например, чтобы доказать первое равенство, нужно добавить XY' к $X'Y'$ и вынести общий множитель.

Реле и переключатели специальных типов

В некоторых типах схем необходимо сохранять определенную зависимость последовательности срабатывания контактов. Это имеет место в случае контактов с размыканием после замыкания и в случае перекидных контактов. Для схем такого рода простейшим методом представляется метод, состоящий в следующем: при составлении уравнений предполагается, что замыкающие и размыкающие контакты срабатывают одновременно, и, после того как произведены все упрощения уравнений и построена соответствующая

схема, нужный тип последовательности действий контактов находится путем дополнительного исследования.

Реле, у которых размыкание и замыкание контактов осуществляется с задержкой по времени, могут трактоваться аналогичным образом или со сдвигом по временной оси. Так если обмотка реле соединена с батареей через сопротивление X и реле имеет задержку на p секунд, то функция сопротивления контактов этого реле также будет иметь значение X , но через p секунд. Это может быть указано введением обозначения $X(t)$ для сопротивления, последовательно соединенного с реле, и обозначения $X(t-p)$ для сопротивления контактов реле.

Существует много типов реле и переключателей, предназначенных для специальных целей; к ним относятся различные шаговые и селекторные переключатели, реле с несколькими обмотками, коммутаторы. Действие реле и переключателей всех этих типов могут быть описаны при помощи слов «или», «и», «если... то...», «срабатывает», «не срабатывает». Это является достаточным условием того, что в терминах функций сопротивления они могут быть описаны операциями сложения, умножения, отрицания и равенства. Так, реле с двумя обмотками может быть сконструировано таким образом, что оно срабатывает, если первая и вторая обмотки возбуждены или если первая и вторая обмотки не возбуждены. Если первая обмотка есть X , а вторая — Y , то функция сопротивления замыкающих контактов этого реле будет $XY + X'Y'$. Обычно, однако, эти специальные реле появляются только на выходе сложных схем и могут быть исключены из рассуждений до тех пор, пока не будет построена остальная часть схемы.

Иногда реле X должно срабатывать, когда схема R замыкается, и оставаться в этом состоянии независимо от состояния схемы R до тех пор, пока не разомкнется схема S . Такие схемы называются схемами с блокировкой. Их уравнение имеет вид

$$X = RX + S.$$

Замена X на X' дает

$$X' = RX' + S$$

или

$$X = (R' + X) S'.$$

В этом случае контакт X размыкается, когда замыкается схема R , и остается разомкнутым до тех пор, пока не разомкнется S .

4. Синтез схем

Некоторые основные теоремы о схемах и функциях

Было показано, что любая функция может быть разложена в сумму произведений, каждое из которых имеет вид $X_1, X_2, \dots, X_n$

с некоторым распределением знаков отрицания у символов и берется с коэффициентом 0 или 1. Так как каждая из n переменных может иметь или не иметь знака отрицания, то имеется 2^n произведений этого вида. Аналогично, так как каждое произведение берется с коэффициентом 0 или 1, то имеется 2^{2^n} возможных сумм такого вида. Следовательно, имеет место

Теорема. Число функций, зависящих от n переменных, равно 2^{2^n} .

Каждая из этих сумм представляет отличную от других функцию, но некоторые из функций могут фактически зависеть менее чем от n переменных (т. е. они имеют такой вид, что для одной или большего числа из n переменных, скажем для X_k , имеем тождественно $f|_{X_k=0} = f|_{X_k=1}$, так что ни при каких условиях значение функции не зависит от значения X_k). Так, для двух переменных X и Y среди 16 функций имеются функции X , Y , X' , Y' , 0 и 1, которые не зависят по крайней мере от одной из переменных X и Y . Чтобы найти число функций, существенно зависящих от всех n переменных, поступим следующим образом. Пусть это число есть $\varphi(n)$. Тогда по только что приведенной теореме

$$2^{2^n} = \sum_{k=0}^n C_n^k \varphi(k),$$

где $C_n^k = n!/k!(n-k)!$ есть число сочетаний из n по k . Это значит, что общее число функций, которые могут быть получены от n переменных, равно сумме чисел таких функций, которые могут быть получены от всех возможных выборов из всех n переменных и которые зависят существенно от всех переменных в данной выборке. Решая последнее уравнение относительно $\varphi(n)$, получим

$$\varphi(n) = 2^{2^n} - \sum_{k=0}^{n-1} C_n^k \varphi(k).$$

Подставляя в правую часть равенства вместо $\varphi(n-1)$ аналогичное выражение, полученное заменой n на $n-1$, затем аналогично — для $\varphi(n-2)$ в полученном выражении и т. д., можно получить уравнение, содержащее только $\varphi(n)$. Это уравнение может быть затем приведено к виду

$$\varphi(n) = \sum_{k=0}^n C_n^k 2^{2^k} (-1)^{n-k}.$$

С ростом n это выражение асимптотически приближается к его главному члену 2^{2^n} . При $n=5$ погрешность при использовании только этого члена не превышает 0,01%.

Построим теперь такие функции от n переменных, которые требуют для своей реализации наибольшего числа контактов, и найдем это число. Для этого необходимо определить функцию двух переменных — знакомую нам сумму по модулю 2. Эта функция обозначается через $X_1 \oplus X_2$ и определяется уравнением

$$X_1 \oplus X_2 = X_1 X_2' + X_1' X_2.$$

Легко показать, что эта операция подчиняется коммутативному и ассоциативному законам и закону дистрибутивности относительно умножения, т. е.

$$\begin{aligned} X_1 \oplus X_2 &= X_2 \oplus X_1; \\ (X_1 \oplus X_2) \oplus X_3 &= X_1 \oplus (X_2 \oplus X_3); \\ X_1 (X_2 \oplus X_3) &= X_1 X_2 \oplus X_1 X_3. \end{aligned}$$

Имеют место также формулы

$$\begin{aligned} (X_1 \oplus X_2)' &= X_1 \oplus X_2' = X_1' \oplus X_2; \\ X_1 \oplus 0 &= X_1; \\ X_1 \oplus 1 &= X_1'. \end{aligned}$$

Так как сумма по модулю два подчиняется ассоциативному закону, можно в сумме нескольких слагаемых опускать скобки. Сумму по модулю два n переменных $X_1, X_2, \dots, X_n$ условимся обозначать через $\bigoplus_{k=1}^n X_k$,

$$\bigoplus_{k=1}^n X_k = X_1 \oplus X_2 \oplus \dots \oplus X_n.$$

Т е о р е м а ¹⁾. Функциями, зависящими от n переменных и требующими для своей параллельно-последовательной реализации наибольшего числа элементов (контактов), являются $\bigoplus_{k=1}^n X_k$ и $(\bigoplus_{k=1}^n X_k)'$, и каждая из них требует $(3 \cdot 2^{n-1} - 2)$ элементов.

Проведем доказательство методом математической индукции. Заметим для начала, что это верно для $n = 2$. Имеется 10 функций, зависящих от двух переменных, а именно $XY, X + Y, X'Y, X' + Y, XY', X + Y', X'Y', X' + Y', XY' + X'Y, XY + X'Y'$. Все они, кроме двух последних, требуют по два элемента для своей реализации; две последние требуют по четыре элемента и являются соответственно функциями $X \oplus Y$ и $(X \oplus Y)'$. Итак, теорема верна для $n = 2$. Предполагая теперь, что она верна для $n - 1$, докажем,

¹⁾ Это утверждение является неверным. См. примечание автора на стр. 82.— Прим. ред.

что она верна для n . Любая функция n переменных может быть разложена по n -й переменной следующим образом:

$$f(X_1, X_2, \dots, X_n) = f(X_1, X_2, \dots, X_{n-1}, 1) + X_n' f(X_1, X_2, \dots, X_{n-1}, 0). \quad (19)$$

Здесь члены $f(X_1, X_2, \dots, X_{n-1}, 1)$ и $f(X_1, X_2, \dots, X_{n-1}, 0)$ суть функции $(n-1)$ переменных, и если они для своей реализации требуют наибольшего числа элементов (в классе функций от $n-1$ переменных), то и f требует наибольшего числа элементов (в классе функций от n переменных), если только нет другого способа записи f , требующего меньшего числа элементов. Мы предположили, что наибольшего числа элементов для $(n-1)$ переменных требуют функции $\bigoplus_{k=1}^{n-1} X_k$ и $(\bigoplus_{k=1}^{n-1} X_k)'$. Если, следовательно, подста-

вить вместо функции $f(X_1, \dots, X_{n-1}, 1)$ функцию $\bigoplus_{k=1}^{n-1} X_k$, а вместо функции $f(X_1, \dots, X_{n-1}, 0)$ — функцию $(\bigoplus_{k=1}^{n-1} X_k)'$, то получим

$$f = X_n (\bigoplus_{k=1}^{n-1} X_k)' + X_n \bigoplus_{k=1}^{n-1} X_k = (\bigoplus_{k=1}^n X_k)'.$$

В силу симметрии не существует другого метода разложения этой функции, снижающего число элементов. Если указанные функции будут подставлены в другом порядке, получим

$$f = X_n (\bigoplus_{k=1}^{n-1} X_k)' + X_n' \bigoplus_{k=1}^{n-1} X_k = \bigoplus_{k=1}^n X_k.$$

Этим завершается доказательство того, что указанные функции требуют наибольшего числа элементов.

Чтобы доказать, что они требуют $3 \cdot 2^{n-1} - 2$ элементов, обозначим число требуемых элементов через $s(n)$. Тогда из равенства (19) получаем разностное уравнение

$$s(n) = 2s(n-1) + 2,$$

где $s(2) = 4$. Оно линейно, имеет постоянные коэффициенты и может быть решено обычными методами. Решением является

$$s(n) = 3 \cdot 2^{n-1} - 2,$$

в чем можно легко убедиться подстановкой в уравнение и проверкой выполнения начальных условий.

Заметим, что вышесказанное применимо только к параллельно-последовательным реализациям. В следующем разделе будет показано, что функция $\bigoplus_{k=1}^n X_k$ и ее отрицание могут быть реализованы

с помощью $4(n-1)$ элементов, если пользоваться более общими типами схем. Функции, требующие наибольшего числа элементов при использовании любых типов схем, пока не построены.

Двойственные схемы

Отрицание (инверсия) произвольной схемы может быть найдено при помощи теоремы де Моргана, но схема должна быть предварительно преобразована в эквивалентную параллельно-последовательную (если она не является схемой этого типа). Докажем теорему,

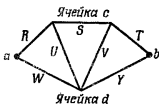


Рис. 15. Плоская схема для иллюстрации теоремы двойственности.

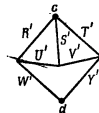


Рис. 16. Схема, двойственная схеме, изображенной на рис. 15.

при помощи которой можно непосредственно найти инверсию любой плоской двухполюсной схемы. В качестве следствия будет дан метод нахождения схемы с постоянным током, эквивалентной данной схеме с постоянным напряжением.

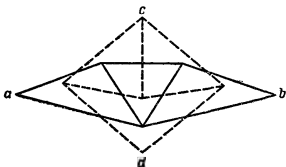
Пусть N — плоская схема с функцией сопротивления X_{ab} между полюсами a и b , находящимися вне схемы. Для определенности рассмотрим схему, изображенную на рис. 15 (где сопротивления изображены просто отрезками).

Пусть теперь M — схема, двойственная схеме N и построенная при помощи следующего процесса: для каждого контура или ячейки схемы N выделен узел схемы M ; каждому элементу X_k схемы N , разделяющему ячейки r и s , поставлен в соответствие элемент X'_k , соединяющий узлы r и s схемы M ; область, внешняя к N , должна рассматриваться как две ячейки, c и d , которым соответствуют полюсы c и d схемы M . Так, схемой, двойственной схеме, изображенной на рис. 15, является схема, изображенная на рис. 16.

Т е о р е м а. Если схема M двойственна схеме N , то $X_{ab} = X'_{cd}$.

Чтобы доказать это, предположим, что M наложена на N так, что узлы схемы M лежат в соответствующих ячейках схемы N , а соответствующие элементы пересекаются. Для схемы, изображенной на рис. 15, это показано на рис. 17: N — сплошной линией, M — пунктиром. Заметим, что самым простым методом нахождения двойственной схемы (как для схем этого типа, так и для схем с ко-

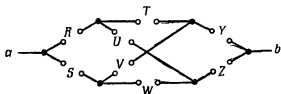
нечными сопротивлениями) является метод, использующий наложение искомой схемы на заданную. Тогда, если $X_{ab} = 0$, то в схеме N должна существовать цепь от a к b такая, что каждый элемент на ней равен нулю. Но эта цепь рассекает схему M между



Р и с. 17. Наложение двойственной схемы на заданную.

полюсами c и d , так что вдоль нее каждый элемент схемы M равен единице. Следовательно, $X_{cd} = 1$. Аналогично, если $X_{cd} = 0$, то $X_{ab} = 1$ и, следовательно, $X_{ab} = X'_{cd}$.

Из этой теоремы, очевидно, вытекает, что инверсия плоской схемы может быть построена с тем же числом элементов, что и заданная схема¹⁾.

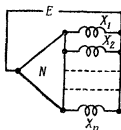


Р и с. 18. Неплоская схема.

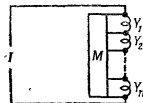
В релейной системе с постоянным напряжением все реле соединены параллельно. Реле срабатывает при замыкании части схемы, включенной с ним последовательно. Общий вид схемы с постоянным напряжением показан на рис. 19. В схеме с постоянным током все реле соединены последовательно. Для отпускания реле его обмотка замыкается накоротко. На рис. 20 изображен общий вид схемы с постоянным током, соответствующей схеме, изображенной на рис. 19. Если реле Y_k на рис. 20 должно срабатывать только тогда, когда сработало реле X_k на рис. 19, то, сопротивление, параллельное Y_k , должно быть инверсным по отношению к сопротивле-

¹⁾ Это, вообще говоря, будет не верно, если опустить слово «плоская». Например, неплоская схема, изображенная на рис. 18, не имеет инверсной схемы, содержащей только 8 элементов.

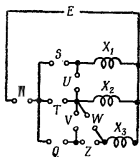
нию X_k , включенному последовательно с обмоткой реле X_k . Если это верно для всех реле, будем говорить, что схемы с постоянным током и постоянным напряжением эквивалентны. Приведенная выше теорема может быть использована для нахождения эквивалентных



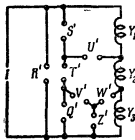
Р и с. 19. Общий вид релейной схемы с постоянным напряжением.



Р и с. 20. Общий вид схемы с постоянным током.



Р и с. 21. Простая система с постоянным напряжением.



Р и с. 22. Система с постоянным током, эквивалентная системе, изображенной на рис. 21.

схем такого рода, так как если построить схемы N и M , двойственные в указанном смысле, в которых элементы X_k и Y_k будут соответствовать друг другу (см. рис. 19 и 20), то условия будут выполнены. Простой пример этого показан на рис. 21 и 22.

Синтез схем, реализующих симметрические функции

Было показано, что любой формуле в точности соответствует параллельно-последовательная схема. Однако параллельно-последовательная реализация может требовать больше элементов, чем некоторые другие схемы, реализующие ту же функцию. В этом разделе будет дан метод построения схем, реализующих функции некоторого

типа, которые, вообще говоря, значительно более экономны по числу элементов, чем наилучшие параллельно-последовательные схемы. Эти функции известны под названием симметрических и часто встречаются в релейных схемах.

О п р е д е л е н и е. Функция n переменных $X_1, X_2, \dots, X_n$ называется симметрической по этим переменным, если перестановка переменных приводит к функции, тождественной данной.

Так, $XY + XZ + YZ$ симметрична по переменным X, Y и Z . Так как любая перестановка переменных может быть получена последовательными перестановками двух переменных, то необходимое и достаточное условие симметричности функций состоит в том, что любая перестановка двух переменных оставляет функцию неизменной.

Надлежащим выбором переменных многие функции, кажущиеся несимметрическими, могут быть сделаны симметрическими. Например, $XY'Z + X'YZ + X'Y'Z'$, хотя и не является симметрической по X, Y и Z , оказывается симметрической по X, Y и Z' . Иногда также возможно несимметрическую функцию представить в виде произведения или суммы симметрической функции и переменной (или переменной с отрицанием). В этом случае симметрическая часть может быть реализована методом, который будет описан, а дополнительный член присоединен параллельно или последовательно.

Этот метод основан на следующей теореме.

Т е о р е м а. Для того чтобы некоторая функция была симметрической, необходимо и достаточно, чтобы ее можно было задать множеством целых чисел $a_1, a_2, \dots, a_n$, таких, что функция обращается в нуль тогда и только тогда, когда число переменных, имеющих значение 0, совпадает с одним из этих чисел.

Это легко вытекает из определения. Множество $a_1, a_2, \dots, a_n$ может быть любым множеством целых чисел от 0 до n , где n — число переменных симметрической функции. Условимся их называть a -числами функции¹⁾. Симметрическая функция $XY + XZ + YZ$ имеет a -числа 2 и 3, так как функция равна нулю только тогда, когда две или три переменные равны нулю. Чтобы найти a -числа данной симметрической функции, необходимо только вычислить значение функции, когда 0, 1, ..., n переменных равны нулю. Те числа, при которых результат есть нуль, являются a -числами функции.

¹⁾ В нашей литературе используется двоякая терминология. Симметрическая функция задается системой рабочих чисел. Рабочее число — это число единиц в наборе, на котором функция обращается в единицу. — *Прим. ред.*

Теорема. *Существует 2^{n+1} симметрических функций от n переменных.*

Это следует из того, что имеется $n + 1$ чисел, каждое из которых может быть взято (или не взято) в качестве a -числа. Однако две из этих функций тривиальны, а именно те, для которых выбраны все числа или не выбрано ни одно число. Это дает функции 0 и 1 соответственно. Симметрическую функцию от n переменных $X_1, X_2, \dots, X_n$ с a -числами $a_1, a_2, \dots, a_k$ будем в дальнейшем обозначать через $S_{a_1, a_2, \dots, a_k}(X_1, X_2, \dots, X_n)$. Например, рассматриваемая нами функция будет обозначаться через $S_{2,3}(X, Y, Z)$. Конструкция схемы, реализующей любую симметрическую функцию, основана на a -числах, и предполагается, что эти числа известны.

Теорема. *Сумма двух симметрических функций от одних и тех же переменных есть симметрическая функция от этих переменных, имеющая в качестве a -чисел те числа, которые являются a -числами для обеих функций.*

Так,

$$S_{1,2,3}(X_1, X_2, \dots, X_6) + S_{2,3,5}(X_1, X_2, \dots, X_6) = S_{2,3}(X_1, X_2, \dots, X_6).$$

Теорема. *Произведение двух симметрических функций от одних и тех же переменных есть симметрическая функция от тех же переменных, имеющая в качестве a -чисел a -числа как одной, так и другой функции.*

Так,

$$S_{1,2,3}(X_1, X_2, \dots, X_6) S_{2,3,5}(X_1, X_2, \dots, X_6) = S_{1,2,3,5}(X_1, X_2, \dots, X_6).$$

Для доказательства этих теорем достаточно заметить, что произведение равно нулю, если один из сомножителей равен нулю, а сумма равна нулю, только если оба члена равны нулю.

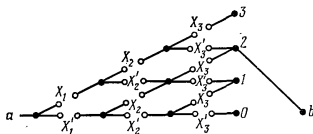
Теорема. *Отрицание симметрической функции от n переменных есть симметрическая функция от этих же переменных, имеющая своими a -числами все числа от 0 до n включительно, не являющиеся a -числами заданной функции.*

Так,

$$S_{2,3,5}(X_1, X_2, \dots, X_6)' = S_{0,1,4,6}(X_1, X_2, \dots, X_6).$$

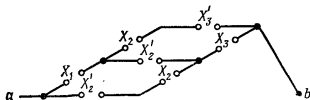
Прежде чем рассмотреть синтез схемы для любой симметрической функции $S_{a_1, a_2, \dots, a_k}(X_1, X_2, \dots, X_n)$, приведем простой пример. Предположим, что надо реализовать функцию $S_2(X_1, X_2, X_3)$. Это значит, что требуется построить схему, которая будет замкнута, если две переменные равны нулю, и разомкнута, если ни одной, одна или три переменных равны нулю. Такая схема изображена на рис. 23. Эта схема может быть разделена на три яруса, каждый от

одной переменной, и четыре уровня, помеченных справа числами 0, 1, 2 и 3. Полюс b соединен с уровнями, соответствующими a -числам реализуемой функции, в данном случае — с уровнем, помеченным цифрой 2. Напряжение подается на полюс a . Если $X_1 = 0$, то напряжение переключается на уровень, помеченный цифрой 1, что соответствует равенству нулю одной переменной. Если же $X_1 = 1$, то напряжение остается на прежнем уровне (полюс a лежит на



Р и с. 23. Схема, реализующая $S_2(X_1, X_2, X_3)$.

нулевом уровне). Затем переходим к переменной X_2 . Если $X_2 = 0$, то напряжение переключается на следующий уровень, если же $X_2 = 1$, то напряжение остается на том же уровне. Таким же образом действует ярус, соответствующий переменной X_3 . Достигнув



Р и с. 24. Упрощение схемы, изображенной на рис. 23.

в конце концов полюсов, расположенных справа, напряжение будет подключено к уровню, номер которого совпадает с общим числом переменных, равных нулю. Так как полюс b соединен с уровнем с пометкой 2, то схема $a-b$ будет замкнута тогда и только тогда, когда 2 из переменных равны нулю. Если бы рассматривалась функция $S_{0,3}(X_1, X_2, X_3)$, полюс b надо было бы соединить с уровнями 0 и 3. На рис. 23 некоторые из элементов, очевидно, лишние. Схема может быть упрощена до вида, изображенного на рис. 24.

Реализация произвольной симметрической функции осуществляется аналогичным образом: в общей схеме для n переменных, изображенной на рис. 25, полюс b соединяют с уровнями, соответствующими a -числам заданной симметрической функции. На рис. 25 сопротивления представлены только отрезками, а буквы опущены,

но сопротивление каждого отрезка можно легко усмотреть по аналогии с рис. 23. После того как полюс b будет присоединен, все лишние элементы могут быть отброшены.

В некоторых случаях возможно значительно упростить схему, совмещая уровни. Пусть задана функция $S_{0,3,6}(X_1, \dots, X_6)$.

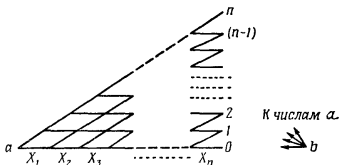


Рис. 25. Схемы реализации симметрической функции $S_{a_1, a_2, \dots, a_k}(X_1, \dots, X_n)$. Сопротивление каждого наклонного элемента равно переменной, написанной под ним; сопротивление каждого горизонтального элемента равно отрицанию этой переменной. Введенные обозначения будут использоваться и далее.

Вместо того чтобы строить схему с шестью уровнями, присоединим второй уровень к нулевому, как показано на рис. 26. На нулевом уровне тогда расположатся также третий и шестой. Присоединяя

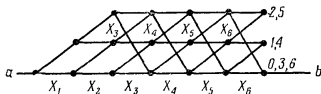
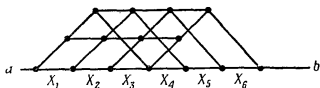


Рис. 26. Схема для $S_{0,3,6}(X_1, X_2, \dots, X_6)$, полученная с использованием процесса совмещения уровней.

полюс b к этому уровню, получим реализацию функции с большой экономией элементов. Отбрасывая элементы, не являющиеся необходимыми, получим схему, изображенную на рис. 27. Этот метод особенно полезен, когда a -числа образуют арифметическую прогрессию, хотя он может иногда применяться и в других случаях.

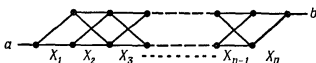
Функции $\sum_{h=1}^n X_h$ и $(\sum_{h=1}^n X_h)'$, которые, как было показано, требуют наибольшего числа элементов при параллельно-последова-

тельной реализации, указанным методом реализуются очень простыми схемами. Можно легко показать, что $\bigoplus_{k=1}^n X_k$ есть симметрическая функция и что, если n четное, α -числами ее являются все



Р и с. 27. Упрощение схемы, изображенной на рис. 26.

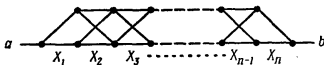
четные числа, а если n нечетно, то — все нечетные числа. Для функции $(\bigoplus_{k=1}^n X_k)'$ имеет место обратное. Используя процесс совмещения, получим схемы, приведенные на рис. 28 и 29. Каждая из



Р и с. 28. Схема для $\bigoplus_{k=1}^n X_k$ при нечетном n

и для $(\bigoplus_{k=1}^n X_k)'$ при четном n .

этих схем требует $4(n-1)$ элементов. В них легко узнать обычную схему включения освещения из n пунктов, использующую $n-2$



Р и с. 29. Схема для $\bigoplus_{k=1}^n X_k$ при четном n

и для $(\bigoplus_{k=1}^n X_k)'$ при нечетном n .

двухполюсных и два однополюсных переключателя на два направления. Если в одной из этих точек меняется положение переключателя, то общее число переменных, равных нулю, изменяется на единицу, так что, если свет включен, он будет выключен, а если нет — включен.

Схемой, изображенной на рис. 25, можно реализовать несколько симметрических функций от одного и того же множества переменных, если только разные функции не имеют общих a -чисел. Если имеются общие a -числа, уровни могут быть совмещены или может быть добавлено дополнительное реле так, что одной схемы будет достаточно.

Универсальная схема, изображенная на рис. 25, содержит $n(n+1)$ элементов. Покажем, что при любом выборе a -чисел по крайней мере n элементов излишни. Каждое число от 1 до $n-1$ включительно, которое не входит в множества a -чисел, порождает два элемента, не являющиеся необходимыми; 0 и n , не являющиеся a -числами, порождают один лишний элемент. Если два a -числа отличаются только на единицу, то два элемента излишни. Если имеется более двух соседних a -чисел или если два или более соседних числа не являются a -числами, то каждое из них порождает более одного лишнего элемента. Тогда очевидно, что наихудшим будет случай, когда a -числами являются все нечетные или все четные числа от 0 до n . В каждом из этих случаев легко видеть, что n элементов будут лишними. В этих случаях процесс совмещения уровней может быть применен, если $n > 2$, так что максимальное число n^2 элементов будет необходимо только для реализации четырех функций X , X' , $X \oplus Y$ и $(X \oplus Y)'$.

Составление уравнений по рабочим характеристикам

В общем случае имеется некоторое множество независимых переменных A, B, C , которые могут быть внешними переключателями или блокировочными реле. Имеется также множество зависимых переменных $x, y, z, \dots$, соответствующих реле, моторам или другим устройствам, управляемым схемой. Требуется найти схему, которая для всех возможных комбинаций значений независимых переменных дает правильные значения всех зависимых переменных. Общий метод решения основывается на следующих правилах:

1. Для каждого дополнительного такта действия последовательной схемы должны быть введены дополнительные зависимые переменные. Так, если требуется построить схему, действующую в три такта, должны быть введены две дополнительные переменные, соответствующие началу двух последних тактов. Эти дополнительные переменные могут представлять контакты шагового переключателя или реле, которые замыкаются последовательно. Аналогично каждая задержка требует новой переменной, соответствующей некоторому реле. Из постановки задачи обычно ясно, какие дополнительные типы реле могут потребоваться для ее решения.

Связь между рабочими характеристиками и уравнениями

Символ	В терминах срабатывания	В терминах несрабатывания
X	Переключатель или реле X срабатывает	Переключатель или реле X не срабатывает
$=$	Если	Если
X'	Переключатель или реле X не срабатывает	Переключатель или реле X срабатывает
$\cdot$	Или	И
$+$	И	Или
$(- -)'$	Схема $(- -)$ не замкнута или трактуется по теореме де Моргана	Схема $(- -)$ замкнута или трактуется по теореме де Моргана
$X(t-p)$	X сработало по крайней мере на p секунд раньше	X было разомкнуто по крайней мере на p секунд раньше

Если зависимые переменные появляются в определяющей их функции (как в схеме с блокировкой), то строгое соблюдение вышесказанного может повести к составлению неправильных уравнений. В таких случаях должны быть использованы следующие зависимости:

$$X = RX + S$$

X срабатывает, когда R замкнуто (при условии, что S замкнуто) и остается замкнутым независимо от состояния R , пока не разомкнется S

$$X = (R' + X) S'$$

X размыкается, когда R замыкается (при условии, что S замкнуто), и остается разомкнутым независимо от состояния R , пока S не разомкнется

При пользовании этой таблицей рекомендуется записывать рассматриваемую функцию либо как сумму только чистых произведений, либо как произведение чистых сумм. В случае суммы произведений условия функционирования должны быть определены в терминах несрабатывания; для произведения сумм — в терминах срабатывания. В тех случаях, когда запись не удовлетворяет этим требованиям, иногда бывает затруднительно расставить скобки в соответствующих выражениях.

2. Для каждой из зависимых переменных должны быть записаны уравнения функции сопротивления. Эти уравнения могут содержать любые переменные, зависимые или независимые, включая переменные, функции которых должны быть определены (как, например, в схемах с блокировкой). Условия могут касаться как срабатывания, так и несрабатывания. Уравнения составляются по условиям функционирования схемы согласно табл. II. Для иллюстрации использования этой таблицы предположим, что реле U должно сработать, если x сработало и y или z сработало, а v , или w , или z не сработало. Выражением для U будет

$$U = x + yz + v'w'z'.$$

Уравнения для реле с блокировкой уже были рассмотрены. Может быть, конечно, что одни и те же условия входят в выражения несколько раз. При окончательном упрощении лишние члены, появляющиеся в таком случае, исчезают.

3. Выражения для различных зависимых переменных должны быть по возможности упрощены при помощи теорем об операциях над этими величинами. Насколько значительным будет это упрощение — до некоторой степени зависит от опыта и изобретательности инженера-конструктора.

4. Наконец должна быть вычерчена окончательная схема. В нее должны быть внесены некоторые необходимые дополнения, продиктованные практическими соображениями, такими, как способность проводить ток, последовательность срабатывания реле и т. п.

5. Примеры

В этой части статьи будет решено несколько задач изложенным методом. Примеры предназначены скорее для того, чтобы проиллюстрировать использование методов исчисления схем в практических задачах и показать разнообразие релейных и переключательных схем, а не для описания фактических устройств.

При помощи релейных схем можно решать сложные математические проблемы. Числа могут быть представлены положениями реле или шаговых переключателей, а для представления различных математических операций могут быть использованы взаимосвязи между множествами реле. Действительно, любая операция, которая может быть полностью описана конечным числом шагов при помощи слов «если ..., то...», «или», «и» и т. д. (см. табл. II), может быть выполнена автоматически посредством реле. Последний пример является иллюстрацией того, как с помощью реле осуществляется одна из математических операций.

Селекторная схема

Реле U должно срабатывать тогда и только тогда, когда срабатывает любое одно, любые три или все четыре реле w , x , y и z . Функция сопротивления U будет, очевидно,

$$U = wxyz + w'x'yz + w'xy'z + w'xyz' + wx'y'z + wx'yz' + wxy'z'$$

или после упрощений, оставаясь в классе параллельно-последовательных схем,

$$U = w[x(yz + y'z') + x'(y'z + yz')] + w'[x(y'z + yz') + x'yz].$$

Такая схема показана на рис. 30. Она требует 20 элементов.

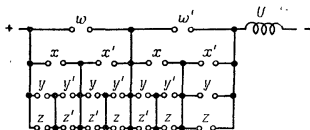


Рис. 30. Параллельно-последовательная реализация селекторной схемы.

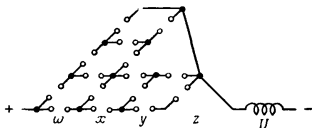


Рис. 31. Селекторная схема, полученная методом реализации симметрических функций.

Используя метод симметрических функций, мы можем написать

$$U = S_{1,3,4}(w, x, y, z).$$

Соответствующая схема (рис. 31) содержит только 15 элементов. Еще большее упрощение может быть получено следующим способом. Сначала рассмотрим

$$U' = S_{0,2}(w, x, y, z).$$

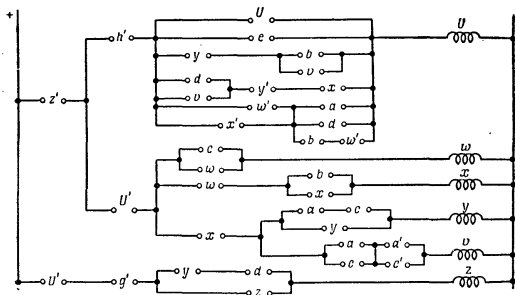


Рис. 34. Схема замка с электрическим секретом. Z и Z' исходно замкнуты, а затем размыкаются. U и U' исходно замкнуты, а затем размыкаются.

довательности нажатия кнопок, обозначим соответственно через w, x, y, z . Из-за необходимости выполнения третьего шага дополнительно требуется реле с временной задержкой. Очевидно, что даже при правильном действии человек не может нажать кнопки a и c в точности одновременно, а как только одна из них нажата, должен сработать сигнал тревоги. Поэтому требуется дополнительное реле времени v , которое будет срабатывать, если после второго шага будет нажата только одна из кнопок a или c дольше, чем на время s задержки реле v .

Если z срабатывает, замок открывается, и в этот момент все остальные реле должны отключиться от схемы. Уравнения системы могут быть записаны непосредственно:

$$w = cw + z' + U';$$

$$x = bx + w + z' + U';$$

$$y = (a + c)y + x + z' + U';$$

$$z = z(d + y) + g' + U';$$

$$v = x + ac + a'c' + z' + U';$$

$$U = e(w' + abd)(w + x' + ad)[x + y' + dv(t - s)] + [y + bv(t - s)]U + h' + z'.$$

Эти выражения могут быть значительно упрощены, сначала совместным преобразованием второго и третьего множителей первого члена из U , а затем выделением общих членов различных функций.

Окончательно упрощенная форма будет

$$\begin{aligned}
 U &= h' + e[ad(b+w') + x'w'](x+y'+dv)(y+vb)U; \\
 \omega &= c\omega; \\
 x &= z' + \left| \begin{array}{l} bx + \omega; \\ U' + \left| \begin{array}{l} (a+c)y; \\ x + \left| \begin{array}{l} ac + a'c'; \\ \end{array} \right. \end{array} \right. \\
 y &= \\
 v &= \\
 z &= g' + \left| (y+d)z + U'.
 \end{aligned}$$

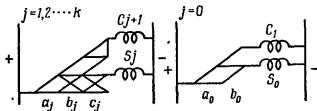
Этим формулам соответствует схема, изображенная на рис. 34.

Двоичный электрический сумматор

Требуется построить схему, которая автоматически складывала бы два числа, используя только реле и переключатели. Хотя может быть использована любая система счисления, наиболее простая схема получается при использовании двоичной системы. Цифрами в ней являются 0 или 1; число, упорядоченные цифры которого суть

$$a_k, a_{k-1}, a_{k-2}, \dots, a_2, a_1, a_0, \text{ есть } \sum_{j=0}^k a_j 2^j.$$

Пусть двум числам, которые надо сложить, соответствуют последовательности переключателей. При этом последовательность



Р и с. 35. Схемы электрического сумматора.

$a_k, a_{k-1}, \dots, a_1, a_0$ представляет цифры одного, а $b_k, b_{k-1}, \dots, b_0$ — цифры другого числа. Сумме будут соответствовать состояния множества реле $s_{k+1}, s, s_{k-1}, \dots, s_1, s_0$. Пусть числу, которое переносится из $(j-1)$ -го разряда в j -й, будет соответствовать реле c_j . Если значение одной из цифр есть 0, соответствующее реле или переключатель будут считаться находящимися в положении, имеющем сопротивление 0; если эта цифра есть 1, реле или переключатель

находится в положении, имеющем сопротивление 1. Сложение показано ниже:

$c_{k+1}, c_k, \dots, c_{j+1}, c_j, \dots, c_2, c_1$	Перенос
$a_k, \dots, a_{j+1}, a_j, \dots, a_2, a_1, a_0$	Первое число
$b_k, \dots, b_{j+1}, b_j, \dots, b_2, b_1, b_0$	Второе число
$c_{k+1}, s_k, \dots, s_{j+1}, s_j, \dots, s_2, s_1, s_0$	Сумма
или	
s_{k+1}	

s_0 равно единице тогда и только тогда, когда $a_0 = 1$ и $b_0 = 0$ или $a_0 = 0$ и $b_0 = 1$.

Следовательно,

$$s_0 = a_0 b'_0 + a'_0 b_0 = a_0 \oplus b_0.$$

$c_1 = 1$ тогда и только тогда, когда $a_0 = 1$ и $b_0 = 1$;

$$c_1 = a_0 b_0.$$

$s_j = 1$, если одна из переменных a_j, b_j, c_j равна единице или если все три равны единице:

$$s_j = S_{1,3}(a_j, b_j, c_j), \quad j = 1, 2, \dots, k.$$

$c_{j+1} = 1$, если две или три из этих переменных равны единице:

$$c_{j+1} = S_{2,3}(a_j, b_j, c_j), \quad j = 1, 2, \dots, k.$$

Используя метод реализации симметрических функций и совмещение уровней по S_j , получим схему, изображенную на рис. 35.

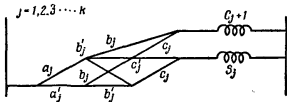


Рис. 36. Упрощение схем, изображенных на рис. 35.

Отбрасывая лишние элементы, приходим к схеме, изображенной на рис. 36.

ЧИСЛО ДВУХПОЛЮСНЫХ ПАРАЛЛЕЛЬНО-ПОСЛЕДОВАТЕЛЬНЫХ СЕТЕЙ¹⁾

Одна из первых попыток описать все электрические схемы, удовлетворяющие некоторым определенным условиям, была сделана в 1892 г. П. А. Мак-Магоном²⁾, который исследовал комбинации сопротивлений при последовательном и параллельном соединении и дал без доказательства формулу производящей функции, по которой может быть определено число таких комбинаций, и таблицу чисел таких комбинаций из 10 или менее элементов³⁾.

Параллельно-последовательные комбинации не исчерпывают всех возможных сетей⁴⁾, так как они исключают мостиковые соединения типа мостика Уитстона⁵⁾, но в силу своей простоты они представляют собой важный подкласс. Если число элементов меньше 5, все сети являются параллельно-последовательными; для 5 имеется одна сеть мостикового типа — мостик Уитстона; при возрастании числа элементов число сетей мостикового типа растет быстрее, чем параллельно-последовательных, так, например, при 9 элементах мостиковые сети составляют около 40% от всех сетей. Из этого можно заключить (и установлено, что это верно), что для большого числа элементов параллельно-последовательные сети составляют относительно малую часть всех сетей. Тем не менее проблема описания всех сетей настолько трудна, что подробное изучение параллельно-

¹⁾ Riordan J., Shannon C., The number of two-terminal series-parallel networks, *Journ. of Math. and Phys.*, 21, 2 (1942), 83.

²⁾ MacMahon P., The combination of resistances, *The Electrician*, April 8, 1892. Cayley A., Collected works, III, 203 (стр. 242 по вопросу о получении производящей функции в других задачах)

³⁾ Здесь можно отметить, что приведенное там число 4984 для 10 элементов неверно, правильное их число, а именно 4624, приведено в табл. I ниже.

⁴⁾ В этом переводе и дальше термин «сеть» используется для обозначения топологического объекта (графа с выделенными терминами), а «схема» — для обозначения сети, ребрам которой приписаны символы переменных. — Прим. ред.

⁵⁾ Полное перечисление всех возможных сетей из n элементов при малом n с различной их классификацией дано в работе Forster R., The geometrical circuits of electrical networks, *Trans. AIEE*, 51 (1932), 309.

последовательных сетей желательно по тем соображениям, что даже слабый свет лучше, чем темнота.

Кроме этого, параллельно-последовательные сети интересны сами по себе в другом отношении, а именно в связи с синтезом релейных и переключательных схем¹⁾, где важно знать, сколько элементов требуется для реализации произвольной переключательной функции $f(x_1, \dots, x_n)$ от n переменных, т. е. число $N(n)$ такое, что каждая из 2^{2^n} различных функций f может быть реализована с N элементами и по крайней мере одна не может быть реализована с меньшим числом элементов. Верхняя оценка для числа двухполюсных схем с B ребрами определяет нижнюю оценку для N , так как число различных схем с N контактами (с учетом приписывания ребрам символов переменных) должно быть не меньше 2^{2^N} ; т. е. должно быть достаточно схем для реализации всех функций. Этот общий факт остается верным, если ограничиться только классом параллельно-последовательных схем, и поскольку переключательные схемы наиболее просто реализуются именно в такой форме, определение необходимого для этого числа элементов представляет непосредственный интерес.

Эти соображения заставили нас разобрать доказательство теоремы Мак-Магона о производящей функции, которое полностью приводится ниже, разработать рекуррентные соотношения и схемы подсчета и с их помощью дополнить таблицу Мак-Магона, исследовать поведение числа параллельно-последовательных сетей при большом числе элементов и, наконец, применить это к переключательным функциям, упомянутым выше. Эти вопросы изложены в отдельных пунктах²⁾.

¹⁾ Shannon C., A symbolic analysis of relay and switching circuits, *Trans. AIEE*, 57 (1938), 713. [См. наст. сборник, стр. 9. — *Прим. ред.*]

²⁾ Понятие параллельно-последовательного соединения элементов настолько очевидно интуитивно, что формальное определение представляется излишним. Однако, поскольку в литературе не имеется никакого определения, можно привести две эквивалентные формулировки:

1. Сеть N называется параллельно-последовательной сетью с полюсами a и b , если через каждый элемент сети N проходит по крайней мере один путь от a к b , не проходящий дважды через одну и ту же вершину, и никакие два пути не проходят через один и тот же элемент в разных направлениях.

2. Сеть называется параллельно-последовательной, если она является параллельным или последовательным соединением двух параллельно-последовательных сетей. Один элемент есть параллельно-последовательная сеть.

Второе определение является индуктивным. Заметим, что оно позволяет определить эквивалентность параллельно-последовательных соединений следующим образом: две параллельно-последовательные сети являются эквивалентными, если они представляют собой параллельное (или последовательное) соединение одних и тех же сетей.

Заметим также, что сеть называется существенно последовательной (существенно параллельной), если она есть последовательное (параллельное) соединение двух параллельно-последовательных сетей.

1. Получение производящей функции

Для одного элемента, очевидно, возможна только одна сеть — сам элемент. Для 2, 3 и 4 элементов на рис. 1 представлены все параллельно-последовательные сети, разделенные на классы существенно последовательных и существенно параллельных сетей по причинам, которые будут объяснены ниже.

Число элементов	Существенно последовательные	Существенно параллельные	Число сетей
2			2
3	 	 	4
4	    	    	10

Р и с. 1.

Заметим, что здесь не различаются сети, эквивалентные относительно перестановок их последовательных и параллельных частей, так как с точки зрения электротехники порядок их взаимного расположения не существен¹⁾.

Эта классификация показывает двойственность: число существенно последовательных сетей равно числу существенно параллельных и каждая сеть взаимно однозначно соответствует своему аналогу. Правило соответствия состоит в том, что существенно последова-

¹⁾ Например, не различаются сети



и



или



и



тельная сеть становится существенно параллельной, если слова «последовательный» и «параллельный» в описании сети поменять местами.

Для характеристики сетей удобно иметь систему символов. Это может быть сделано следующим образом: использованием знака $+$ для обозначения последовательного соединения; точки или простого присоединения для обозначения параллельного соединения элементов; 1 для обозначения единичного элемента сети и введения сокращений: n для $1 + 1 + \dots + 1$ (n элементов, соединенных последовательно) и 1^n для $1 \cdot 1 \cdot 1 \dots 1$ (n элементов, соединенных параллельно); например, символ 21 изображает параллельное соединение одного элемента с двумя последовательно соединенными элементами.

Тогда сетям, изображенным на рис. 1, соответствуют обозначения, приведенные в следующей таблице:

Число элементов	Существенно последовательные	Существенно параллельные	Число сетей
2	2	1^2	2
3	3	1^3	4
4	$1^2 + 1$	21	10
	4	1^4	
	$1^2 + 2$	21^2	
	$21 + 1$	$(1^2 + 1) 1$	
	$1^3 + 1$	31	
	$1^2 + 1^2$	22	

Рассматривая существенно параллельные сети, можно заметить, что для $n = 2$ и 3 цифровыми обозначениями являются разложения числа n на слагаемые, исключая само n . Если в качестве разложения взято само n для обозначения существенно последовательной сети, то все параллельно-последовательные сети представляются разложениями числа n для $n < 4$. При $n = 4$ появляется обозначение $(1^2 + 1) 1$ (и соответствующая сеть), которое не является разложением числа n . Но $(1^2 + 1)$ есть одна из существенно последовательных сетей для $n = 3$. Отсюда все сети представляются в виде разложений, если каждая часть разложения интерпретируется как вся совокупность соответствующих существенно последовательных сетей, например разложение 31 интерпретируется как сети 31 и $(1^2 + 1)1$.

При подсчете это означает, что каждое разложение имеет относящийся к нему численный коэффициент, определяемый числом существенно последовательных сетей, соответствующих каждой из его составных частей. Если число таких сетей из p элементов обоз-

начается через a_p , то коэффициент для разложения $(pqr\dots)$, где никакие части не повторяются, равен $a_p a_q a_r \dots$ при $a_1 = a_2 = 1$, так как каждая из комбинаций, соответствующих данной части, может быть соединена параллельно с комбинациями, соответствующими остальным частям. Коэффициент для повторяющейся части, скажем p^π (p , повторенное π раз), есть число сочетаний элементов с неограниченными повторениями из a_p по π , т. е. биномиальный коэффициент $C_{a_p + \pi - 1}^\pi$.

Отсюда общее число параллельно-последовательных сетей s_n из n элементов может быть записано в виде

$$s_n = 2a_n = \sum C_{a_p + \pi_1 - 1}^{\pi_1} C_{a_q + \pi_2 - 1}^{\pi_2} \dots, \quad (1)$$

где сумма берется по всем системам неотрицательных целых чисел $p, q, \dots, \pi_1, \pi_2, \dots$ таким, что $p\pi_1 + q\pi_2 + r\pi_3 + \dots = n$ и $a_1 = a_2 = 1$. Другими словами, эта сумма берется по всем разложениям числа n .

Таким образом, для $n = 5$ разложениями являются

$$5, 41, 32, 31^2, 2^21, 21^3, 1^5$$

и

$$s_5 = a_5 + a_4 + a_3 + a_3 + 1 + 1 + 1$$

или, так как $s_n = 2a_n$,

$$s_5 = s_4 + 2s_3 + 6 = 24.$$

Аналогично

$$s_8 = s_5 + 2s_4 + 2C_{a_3+1}^2 + 2s_3 + 8 = 66.$$

Производящая функция¹⁾, данная Мак-Магоном, а именно

$$\prod_{i=1}^{\infty} (1 - x^i)^{-a_i} = 1 + \sum_{n=1}^{\infty} s_n x^n, \quad (2)$$

где $\prod$ означает произведение, может быть получена из (1) рассуждениями, несущественно отличающимися от тех, которые были использованы для получения производящей функции Эйлера²⁾ для разложений числа n , которая имеет вид

$$\prod_{i=1}^{\infty} (1 - x^i)^{-1} = 1 + \sum_{n=1}^{\infty} p_n x^n.$$

¹⁾ Следует заметить, что это не есть производящая функция в том смысле, что коэффициенты при степенях в ряде полностью определяются по формуле, а скорее есть производящее тождество, определяющее коэффициенты путем сравнения членов с одинаковыми степенями.

²⁾ Ср., например, с работой Hardy G., Wright E., An introduction to the theory of numbers, Oxford, 1938, 272.

2. Численный подсчет

Прямое вычисление по производящему тождеству (2) или по эквивалентному ему уравнению (1) становится громоздким даже при относительно малых значениях n , так как число членов равно числу разложений. Более того, вычисление ведется последовательно, каждое число зависит от предшествующих ему и ошибка накапливается; поэтому желательно иметь независимые схемы вычисления.

Ниже приводятся три схемы, используемые при подсчете числа параллельно-последовательных сетей, приведенных в табл. I.

Таблица I

Число параллельно-последовательных сетей и числа σ_n

n	s_n число существ. венно параллельных (или последователь- ных) сетей	числа σ_n	n	s_n число существ. венно параллельных (или последователь- ных) сетей	числа σ_n
1	1	1	16	4 507 352	1 104 347
2	2	1	17	14 611 576	3 584 649
3	4	1	18	47 633 486	11 701 369
4	10	3	19	156 047 204	38 347 065
5	24	5	20	513 477 502	126 395 250
6	66	17	21	1 696 305 720	
7	180	41	22	5 623 993 944	
8	522	127	23	18 706 733 128	
9	1 532	365	24	62 408 176 762	
10	4 624	1 119	25	208 769 240 140	
11	14 136	3 413	26	700 129 713 630	
12	43 930	10 685	27	2 353 386 723 912	
13	137 908	33 561	28	7 927 504 004 640	
14	437 502	106 827	29	26 757 247 573 360	
15	1 399 068	342 129	30	90 479 177 302 242	

Эти схемы тесно связаны со схемами подсчета числа разложений, принадлежащими Эйлеру и Гупта, что неявно выражено в рекуррентных формулах.

Первая схема подсчета существенно зависит от вычисления «близкого» множества чисел $s_n(k)$, определяемых соотношением

$$\prod_{i=1}^k (1 - x^i)^{-a_i} = 1 + \sum_{n=1}^{\infty} s_n(k) x^n, \quad (3)$$

где $s_n = s_n(N)$, $N \geq n$.

Рекуррентная формула для этих чисел следует непосредственно из определения и имеет следующий вид:

$$s_n(k) = \sum_{i=0}^q C_{a_k + i - 1}^i s_{n-ik}(k-1),$$

где q есть целая часть от n/k и $s_0(k-1) = s_0(k) = 1$. Ясно, что $s_n(1) = 1$, $s_n(2) = 1 + \left[\frac{n}{2} \right]$, где квадратные скобки означают «целая часть от».

Заметим, что числа $s_n(k)$ выражают число существенно параллельных (или существенно последовательных) сетей с n элементами, которые могут быть составлены из частей, содержащих каждая не более k элементов; например, существенно параллельные сети, представляемые числом $s_4(2)$, суть 2^2 , 21^2 и 1^4 . Это замечание вместе с интерпретацией биномиальных коэффициентов, приведенной в разд. 1, дает интерпретацию рекуррентного соотношения (4) для сетей.

Как было указано, числа $s_n(k)$ могут быть использованы для подсчета s_n непосредственно, однако они более эффективно используются в следующей формуле:

$$s_n = s_{n-1} + s_{n-2}s_2 + \dots + s_{n-m-1}s_{m+1} + 2s_n(m), \quad (5)$$

где $m = \left[\frac{n}{2} \right]$.

Интерпретация этого равенства для сетей лучше видна из эквивалентной формулы

$$s_n = a_n + a_{n-1}a_1 + a_{n-2}s_2 + \dots + a_{n-m-1}s_{m+1} + s_n(m). \quad (5')$$

Таким образом, полное число сетей с n элементами складывается из a_n существенно последовательных сетей с n элементами, числа существенно параллельных сетей, построенных путем составления всех существенно последовательных сетей из $n-i$ элементов со всеми сетями из i элементов (i изменяется от 1 до наименьшего из чисел $m+1$ и $n-m-1$) и $s_n(m)$ сетей, описанных выше.

Это в основном все, что используется в так называемом подсчете по Эйлеру.

Подсчет по Гупта основывается на разделении разложений на классы, соответствующие величине наименьшей части; например, если разложение числа n с наименьшей частью k обозначить через $p_{n,k}$, тогда классами для $n=4$ будут

$$p_{4,1} = (31, 21^2, 1^4),$$

$$p_{4,2} = (2^2),$$

$$p_{4,3} = \text{нет},$$

$$p_{4,4} = (4).$$

Рекуррентные формулы для числа сетей соответствующих классов $s_{n,k}$ получаются подходящей модификацией процесса, данного Гупта; так, например, если единица вычеркнута из всех разложений в $p_{n,1}$, то результат есть в точности p_{n-1} . Отсюда

$$s_{n,1} = s_{n-1}.$$

Подобным образом

$$s_{n,2} = a_2 [s_{n-2,2} + s_{n-2,3} + \dots + s_{n-2,n-2}] = s_{n-2} - s_{n-2,1} = s_{n-2} - s_{n-3}.$$

Вообще

$$s_{n,k} = \sum_{i=1}^q C_{a_k+i-1}^i A_{n-ik,k}, \quad (6)$$

где

$$q = \left[\frac{n}{k} \right]$$

$$A_{0,k} = 1,$$

$$A_{r,k} = 0, \quad r = 1, 2, \dots, k,$$

$$A_{r,k} = s_r - s_{r,1} - \dots - s_{r,k}, \quad r > k.$$

Другая форма соотношения (6), получаемая итерацией и более простая чем (6), для малых значений k и больших значений n , имеет вид

$$s_{n,k} = \sum_{i=1}^q C_{a_k}^i A_{n-ik,k-1}. \quad (6')$$

Надо отметить, что в этой сумме появляются пустые члены, если $q > a_k$.

Третья схема подсчета состоит в определении третьего множества чисел σ_n , определяемых так:

$$\prod_{i=1}^{\infty} (1 - x^i)^{a_i} = 1 - \sum_{n=1}^{\infty} \sigma_n x^n. \quad (7)$$

Сопоставляя это определение с производящей функцией Мак-Магона и с равенством (2), получаем

$$s_n = \sum_{i=1}^n \sigma_i s_{n-i}, \quad (8)$$

где s_0 по определению равно единице.

Рекуррентная формула для этих чисел имеет вид

$$\sigma_n = a_n - \sum_{i=1}^{n-m-1} \sigma_i a_{n-i} + \sigma_n(m), \quad (9)$$

где, как и выше, $m = \left[\frac{n}{2} \right]$ и $\sigma_n(k)$ определяются аналогично $s_n(k)$. Заметим, что $\sigma_1 = \sigma_2 = \sigma_3 = 1$. Эти числа включены в табл. I ($n < 20$).

3. Асимптотическое поведение

Поведение s_n для больших n в идеальном случае должно выражаться точной формулой или, если такую формулу не удастся найти, асимптотической формулой. Замечателен тот факт, что асимптотическая формула для числа разложений является «точной» формулой, т. е. может быть использована в вычислении точных значений для больших n . Такие формулы для s_n не найдены; вместо этого даются функции, оценивающие s_n снизу и сверху.

Очевидно прежде всего, что $s_n \geq p_n$ для всех значений n . Но этого мало. Несколько лучшей оценкой является

$$s_n \geq \pi_n, \quad (10)$$

где $\pi_n = 2^{n-1}$ есть число композиций числа n , т. е. разложений числа n , в которых порядок расположения частей существен.

Это доказывается так. Из уравнения (5) имеем $s_n \geq q_n$, если

$$q_n \leq s_n \text{ при } n \leq 4,$$

$$q_n = q_{n-1} + s_2 q_{n-2} \text{ при } n > 4.$$

Решение последнего уравнения, если положить $q_3 = 4, q_4 = 8$, выражается формулой

$$q_n = 2^{n-1} = \pi_n.$$

Если учитывать большее число членов уравнения (5), то нижняя оценка повышается, но относительно медленно, а анализ быстро усложняется; лучшее, что было получено в этом направлении, есть оценка

$$s_n \geq A 3^n, \quad (11)$$

где A — фиксированная константа.

Другой, более интуитивный путь гораздо лучше. Во-первых, заметим, что сетей с n элементами заведомо больше, чем их получается при параллельном или последовательном присоединении одного элемента к сетям с $n-1$ элементом, которые получаются при параллельном или последовательном присоединении одного элемента к сетям с $n-2$ элементами и т. д.¹⁾ Отсюда

$$s_n \geq \pi_n,$$

¹⁾ То есть здесь рассматриваются сети вида



где

$$\pi_n = 2\pi_{n-1} = 2^2\pi_{n-2} = \dots = 2^{n-1}\pi_1 = 2^{n-1},$$

т. е. результат, полученный выше.

Сети из n элементов с одним элементом, присоединенным параллельно (или последовательно), — это как раз те сети, которые представляются числом $s_{n,1}$ в классификации Гупта. Поэтому аппроксимация может быть улучшена рассмотрением большего числа членов в выражении

$$s_n = 2 \sum_{i=1}^m s_{n,i}, \quad m = \left[\frac{n}{2} \right].$$

Член $s_{n,i}$ учитывает существенно последовательные сети, в которых наименьшая существенно параллельная часть имеет точно i элементов. Если из каждой такой сети удалить эту часть, то оставшихся сетей будет не менее, чем существенно последовательных сетей с $n-i$ элементами, если $i < m$; таким образом:

$$s_{n,i} \geq a_i a_{n-i}, \quad i < m.$$

Для n четного, скажем $2m$,

$$s_{2m,m} = C_{a_{m+1}}^a = \frac{1}{2} (a_m^2 + a_m) = \frac{1}{8} (s_m^2 + 2s_m);$$

для n нечетного

$$s_{2m+1,m} = a_{m+1} a_m = \frac{1}{4} s_{m+1} s_m.$$

Отсюда

$$s_{2m} \geq 2s_{n-1} + \frac{1}{2} \sum_{i=2}^{m-1} s_i s_{n-i} + \frac{1}{4} (s_m^2 + 2s_m); \quad (12)$$

$$s_{2m+1} \geq 2s_{n-1} + \frac{1}{2} \sum_{i=2}^m s_i s_{n-i}.$$

Поэтому вообще $s_n \geq r_n$, если $r_1 = 1$, $r_2 = 2$ и

$$r_n = \frac{3}{2} r_{n-1} + \frac{1}{4} \sum_{i=1}^{n-1} r_i r_{n-i}, \quad n > 2. \quad (13)$$

Записывая производящую функцию для r_n в виде

$$R(x) = \sum_{n=1}^{\infty} r_n x^n,$$

из рекуррентного соотношения (13) вместе с начальными условиями получаем

$$[R(x)]^2 - (4 - 6x)R(x) + 4x - x^2 = 0,$$

т. е.

$$R(x) = 2 - 3x - 2\sqrt{1 - 4x + 2x^2}. \quad (14)$$

Асимптотическое поведение r_n может быть определено из $R(x)$ методом Дарбу¹⁾; именно

$$r_n \sim A\lambda^n n^{-3/2}, \quad (15)$$

где A — фиксированная константа и $\lambda = 2 + \sqrt{2} = 3,414 \dots$.

Верхняя оценка получается тем же процессом, если учесть, что

$$s_{n,i} \leq a_i s_{n-i}.$$

Отсюда

$$s_n \leq t_n, \text{ если } t_1 = 1, t_2 = 2$$

и

$$t_n = t_{n-1} + \frac{1}{2} \sum_{i=1}^{n-1} t_i t_{n-i}. \quad (16)$$

Аналогично вышеизложенному

$$T(x) = \sum_{n=0}^{\infty} t_n x^n = 1 - x - \sqrt{1 - 4x} \quad (17)$$

и

$$\begin{aligned} t_n &= \frac{4(2n-3)!}{n!(n-2)!} \approx \\ &\approx \frac{2}{\sqrt{\pi}} 4^{n-1} n^{-3/2}, \quad n > 1 \end{aligned}$$

Сравнение r_n , s_n и t_n для $n \leq 10$, где для удобства берется только целая часть от r_n (обозначается через $[r_n]$), приводится в следующей таблице:

n	1	2	3	4	5	6	7	8	9	10
$[r_n]$	1	2	4	9	22	57	154	429	1225	3565
s_n	1	2	4	10	24	66	180	522	1532	4624
t_n	1	2	4	10	28	84	264	858	2860	9724

¹⁾ Курант Р., Гильберт Д., Методы математической физики, т. II. изд. 2, М., 1951.

Отметим, что для больших значений n нижняя оценка ближе к истинному значению.

Третья, полуэмпирическая оценка ничего не дает. Из табл. I замечаем, что для $n > 2$ число $4s_n$ приближенно равно s_n . Принимая это в качестве равенства и используя уравнение (8) и известные значения $\sigma_1 = \sigma_2 = 1$, получаем уравнение для производящей функции $U(x)$ приближенных значений u_n в виде

$$U(x) = \frac{1}{2} [5 - 3x - 2x^2 - \sqrt{9 - 30x - 11x^2 + 12x^3 + 4x^4}] . \quad (18)$$

Сравнение s_n и целой части u_n показано в таблице II для $n \leq 20$. Приближение в высшей степени точно; наибольшее расхождение

Таблица II

Аппроксимация числа параллельно-последовательных сетей

n	$[u_n]$	s_n	n	$[u_n]$	s_n
1	1	1	11	14 230	14 136
2	2	2	12	44 357	43 930
3	4	4	13	139 779	137 908
4	9	10	14	444 558	437 502
5	23	24	15	1 425 151	1 399 068
6	63	66	16	4 600 339	4 507 352
7	177	180	17	14 939 849	14 611 576
8	514	522	18	48 778 197	47 633 486
9	1527	1532	19	160 019 885	156 047 204
10	4625	4624	20	527 200 711	513 477 502

составляет 10% для $n = 4$, но от $n = 7$ до 20 значения совпадают с точностью до 3%.

Асимптотическое поведение u_n выражается формулой

$$u_n \sim A \lambda^n n^{-3/2},$$

где A около $3/7$, λ около 3,56.

4. Параллельно-последовательная реализация переключательных функций

В качестве приложения полученных результатов покажем, что почти все переключательные функции от n переменных требуют не меньше

$$(1 - \varepsilon) \frac{2^n}{\log_2 n}, \quad \varepsilon > 0,$$

переключающих элементов (замыкающих или размыкающих контактов) для их реализации в классе параллельно-последовательных схем¹⁾.

Число функций, которые могут быть реализованы с h элементами, заведомо не больше числа s_h параллельно-последовательных схем, умноженного на число различных способов, какими могут быть обозначены элементы в каждой схеме. Последнее число равно $(2n)^h$, так как каждый элемент имеет выбор из $2n$ обозначений соответственно каждому переменному и его отрицанию. Поэтому с h элементами могут быть реализованы не более

$$(2n)^h s_h \leq (2n)^h 4^h = (8n)^h$$

различных функций. Если

$$h = \frac{2^n}{\log_2 n} (1 - \varepsilon), \quad \varepsilon > 0,$$

то доля всех 2^{2^n} функций от n переменных, которая может быть реализована, не больше чем

$$\frac{(8n)^{\frac{2^n}{\log_2 n} (1 - \varepsilon)}}{2^{2^n}} = \frac{2^{3(1 - \varepsilon) 2^n \log_n 2 + (1 - \varepsilon) 2^n}}{2^{2^n}} < 2^{3 \cdot 2^n \log_n 2 - \varepsilon \cdot 2^n},$$

и так как последнее выражение стремится к нулю при $n \rightarrow \infty$ для всех положительных ε , то наше утверждение доказано.

¹⁾ Обобщение этой теоремы при весьма общих предположениях получено Р. Е. Кричевским (К р и ч е в с к и й Р. Е., О реализации функций суперпозициями, Сб. «Проблемы кибернетики», вып. 2, Физматгиз, 1959, 123).—
Прим. ред.

СИНТЕЗ ДВУХПОЛЮСНЫХ ПЕРЕКЛЮЧАТЕЛЬНЫХ СХЕМ¹⁾

Часть I

ОБЩАЯ ТЕОРИЯ

1. Введение

Теория переключательных схем может быть разделена на две основные части: анализ и синтез. Задача анализа — определение способа функционирования данной переключательной схемы — относительно проста. Обратная же задача — нахождение схемы, удовлетворяющей заданным условиям функционирования, и, в частности, наилучшей схемы, — задача, вообще говоря, более трудная и более важная с практической точки зрения. Основная часть общей задачи синтеза — это построение двухполюсных схем с данными условиями функционирования, и мы рассмотрим здесь некоторые аспекты этого вопроса.

Переключательные схемы могут быть изучены при помощи булевой алгебры²⁾. Это отрасль математики, впервые исследованная Джорджем Булем в связи с изучением логики и с тех пор применяемая в различных областях, таких, как аксиоматическое построение биологии³⁾, изучение нейронных сетей в нервной системе⁴⁾, анализ страховых полисов⁵⁾, теория вероятностей, теория множеств и т. д.

Возможно, простейшая интерпретация булевой алгебры и одно из ближайших приложений к переключательным схемам — это интерпретация в терминах высказываний. Буква X , например, в этой алгебре соответствует некоторому логическому высказыванию. Сумма двух букв $X + Y$ изображает высказывание « X или Y », а произведение XY изображает высказывание « X и Y ». Символ X'

¹⁾ Shannon C., The synthesis of two-terminal switching circuits, *B. S. T. J.*, 28, № 1 (1949), 59.

²⁾ Birkhoff G., and Mac Lane S., *A survey of modern algebra*, Macmillan, 1949. Couturat L., *The algebra of logic*, Open Court, 1914.

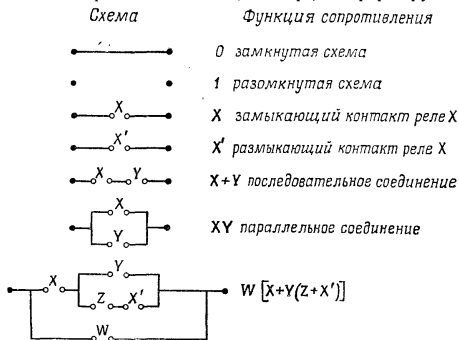
³⁾ Woodger J., *The axiomatic method in biology*, Cambridge; 1937.

⁴⁾ McCulloch W., Pitts W., *A logical calculus of the ideas immanent in nervous activity*, *Bull. Math. Biophysics*, 5 (1943), 115. [Русский перевод: Мак-Каллок У. С. и Питтс Х., *Логическое исчисление идей, относящихся к нервной активности*, сб. Автоматы, М., ИЛ, 1956, 369. — Прим. ред].

⁵⁾ Berkeley E., *Boolean algebra and applications to insurance*, *Record (American Institute of Actuaries)*, 26 (1947), 373.

используется для изображения отрицания высказывания X , т. е. высказывания «не X ». Константы 1 и 0 изображают соответственно истинность и ложность. Так, $X + Y = 1$ означает, что X или Y истинно, а $X + YZ' = 0$ означает, что X или (Y и отрицание Z) ложно.

Интерпретация булевой алгебры в терминах переключательных схем¹⁾ очень проста. Символ X (в алгебре) интерпретируется как



Р и с. 1. Функция сопротивления для простых схем.

замыкающий (front) контакт реле или переключателя. Отрицание X (пишется X') интерпретируется как размыкающий (back) контакт того же реле или переключателя. Константы 0 и 1 интерпретируются соответственно замкнутыми и разомкнутыми схемами, а операции сложения и умножения — последовательными и параллельными соединениями соответствующих переключательных элементов, как это показано на рис. 1. При помощи этих обозначений можно написать алгебраическое выражение, соответствующее двухполюсной схеме. Это выражение, содержащее наименования различных реле, контакты которых встречаются в схеме, будем называть сопротив-

¹⁾ Shannon C., A symbolic analysis of relay and switching circuits, *Trans. AIEE*, 57 (1938), 713. [См. стр. 9 данной книги. — Прим. ред.]; Nakashima A., различные статьи в *Nippon electrical communication engineering*, April, Sept., Nov., Dec., 1938. Piesch H., статьи из архива *Electrotechnik*, 33 (1939), 692, 733. Montgomerie G., Sketch for an algebra of relay and contact circuits, *Journ. of E. E.*, 95, III, № 36 July (1948), 303.

лением или функцией сопротивления схемы. Простым примером является последняя схема на рис. 1.

С булевыми выражениями можно обращаться так же, как с обычными алгебраическими выражениями. Их члены можно располагать в другом порядке, перемножать, умножать на постоянные коэффициенты и комбинировать согласно всем правилам численной алгебры. Так, например, в булевой алгебре имеем следующие тождества:

$$0 + X = X$$

$$0 \cdot X = 0$$

$$1 \cdot X = X$$

$$X + Y = Y + X$$

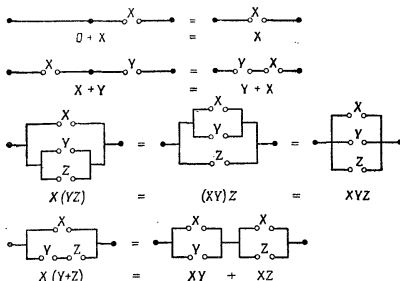
$$XY = YX$$

$$X + (Y + Z) = (X + Y) + Z$$

$$X(YZ) = (XY)Z$$

$$X(Y + Z) = XY + XZ.$$

Интерпретация некоторых из них в терминах переключательных схем показана на рис. 2.



Р и с. 2. Интерпретация некоторых алгебраических тождеств.

Имеется также несколько особых правил булевой алгебры, которые позволяют производить упрощения выражений, невозможные в обычной алгебре. Наиболее важными из них являются:

$$X = X + X = X + X + X \text{ и т. д.}$$

$$X = X \cdot X = X \cdot X \cdot X \text{ и т. д.}$$

$$X + YZ = (X + Y)(X + Z)$$

$$X + 1 = 1$$

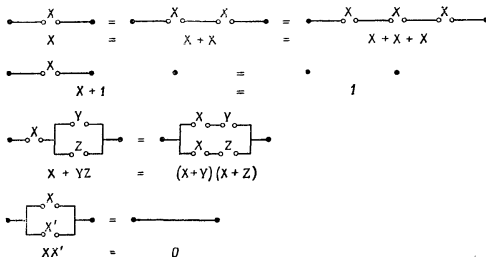
$$X + X' = 1$$

$$X \cdot X' = 0$$

$$(X + Y)' = X' \cdot Y'$$

$$(XY)' = X' + Y'$$

Схемная интерпретация некоторых из них показана на рис. 3. Благодаря этим правилам операции над булевыми выражениями оказываются значительно проще, чем в обычной алгебре. Отпадает, например, необходимость в числовых коэффициентах или в показателях степени, так как $nX = X^n = X$.



Р и с. 3. Интерпретация некоторых специальных тождеств булевой алгебры.

Средствами булевой алгебры можно найти много схем, функционально эквивалентных данной схеме. Можно записать сопротивление данной схемы и преобразовывать его согласно указанным правилам. Каждое новое выражение, полученное в результате этих операций, представляет новую схему, эквивалентную данной. В частности, выражения можно преобразовать таким образом, что не являющиеся необходимыми элементы будут исключены, что приведет к упрощению схемы.

Произвольное выражение, содержащее некоторое число переменных $X_1, X_2, \dots, X_n$, называется *функцией* этих переменных и записывается в обычных функциональных обозначениях символом $f(X_1, X_2, \dots, X_n)$, например, $f(X, Y, Z) = X + Y'Z + XZ'$.

В булевой алгебре есть несколько важных общих теорем, справедливых для любых функций. Функцию можно *разложить* по одному или нескольким ее аргументам следующим образом:

$$f(X_1, X_2, \dots, X_n) = X_1 f(1, X_2, \dots, X_n) + X'_1 f(0, X_2, \dots, X_n).$$

Это — разложение по X_1 , где член $f(1, X_2, \dots, X_n)$ — есть функция $f(X_1, X_2, \dots, X_n)$, в которой вместо X_1 подставлена единица, а член $f(0, X_2, \dots, X_n)$ — та же функция, где вместо X_1 подставлен нуль. Разложение по X_1 и X_2 будет иметь вид

$$\begin{aligned} f(X_1, X_2, \dots, X_n) = & X_1 X_2 f(1, 1, X_3, \dots, X_n) + \\ & + X_1 X'_2 f(1, 0, X_3, \dots, X_n) + X'_1 X_2 f(0, 1, X_3, \dots, X_n) + \\ & + X'_1 X'_2 f(0, 0, X_3, \dots, X_n). \end{aligned}$$

Этот процесс можно продолжить до получения разложения по любому числу переменных. Когда это разложение произведено для всех n переменных, f записывается как сумма 2^n произведений, каждое с коэффициентом, не зависящим ни от одного из переменных. Каждый коэффициент, следовательно, есть константа, либо 0, либо 1.

Подобным же образом можно разложить функцию f в произведение

$$\begin{aligned} f(X_1, X_2, \dots, X_n) = & \\ = [X_1 + f(0, X_2, \dots, X_n)] [X'_1 + f(1, X_2, \dots, X_n)] = & \\ = [X_1 + X_2 + f(0, 0, X_3, \dots, X_n)] [X'_1 + X'_2 + f(0, 1, X_3, \dots, X_n)] \cdot & \\ \cdot [X'_1 + X_2 + f(1, 0, X_3, \dots, X_n)] \cdot & \\ \cdot [X'_1 + X'_2 + f(1, 1, X_3, \dots, X_n)] = \text{и т. д.} & \end{aligned}$$

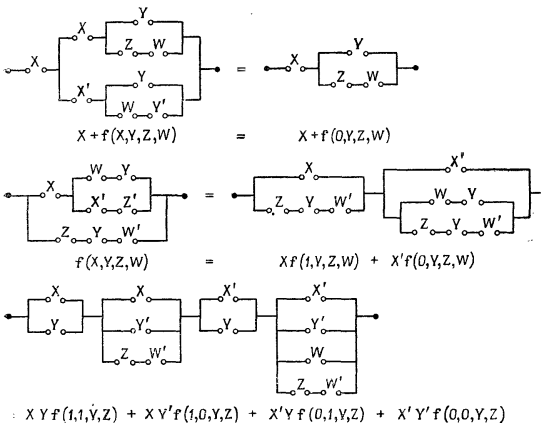
Следующие тождества справедливы для произвольных функций

$$\begin{aligned} X + f(X, Y, Z, \dots) &= X + f(0, Y, Z, \dots), \\ X' + f(X, Y, Z, \dots) &= X' + f(1, Y, Z, \dots), \\ Xf(X, Y, Z, \dots) &= Xf(1, Y, Z, \dots), \\ X'f(X, Y, Z, \dots) &= X'f(0, Y, Z, \dots). \end{aligned}$$

Интерпретация некоторых из этих тождеств показана на рис. 4. Нетрудно заметить, что они справедливы для произвольных переключательных схем.

Функция сопротивления, соответствующая двухполюсной схеме, полностью описывает схему с внешней точки зрения. Можно определить, какая из двух схем разомкнута или замкнута при любом конкретном состоянии реле. Это осуществляется посредством при-

писывания переменным, соответствующим возбужденным реле, значения 0 (так как тогда их замыкающие контакты замкнуты и размыкающие контакты разомкнуты), а переменным, соответствующим не



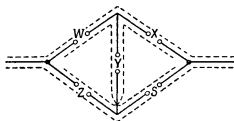
Р и с. 4. Примеры интерпретации некоторых функциональных зависимостей

возбужденным реле, — значения 1. Например, если для функции $f = W[X + Y(Z + X')]$ предположить, что X и Y — возбужденные реле, а Z и W — невозбужденные, то $f = 1[0 + 0(1 + 1)] = 0$, и при этих условиях схема замкнута.

Функция сопротивления в точности соответствует параллельно-последовательному типу схемы, т. е. схеме, содержащей только последовательные и параллельные соединения. Это происходит по той причине, что выражение состоит лишь из операций сложения и умножения. Однако функция сопротивления, представляющая условия функционирования (условия того, замкнута или разомкнута схема между двумя полюсами), существует для схемы любого типа, а не только параллельно-последовательной. Сопротивление схем, не являющихся параллельно-последовательными, может быть най-

дено различными способами, один из которых приведен на рис. 5 для простого мостика. Сопротивление записывается в виде произведения нескольких сомножителей. Каждый из них есть сопротивление возможного пути между двумя полюсами. Дальнейшие детали, относящиеся к булеву методу, применительно к переключательным схемам можно найти в литературе, на которую мы ссылались выше.

Данная статья посвящена проблеме синтеза двухполюсной схемы, реализующей заданную функцию $f(X_1, X_2, \dots, X_n)$. Для любой



$$f = (W+X)(Z+S)(W+Y+S)(Z+Y+X)$$

Р и с. 5. Сопротивление мостиковой схемы.

заданной функции f имеется неограниченное число схемных реализаций, и выбор в каждом случае может быть продиктован различными соображениями. Наиболее обычными являются те или иные ограничения, накладываемые на сложность схемы, например:

1. Реализовать нашу функцию схемой с наименьшим общим числом переключающих элементов независимо от того, какие переменные они представляют.

2. Найти схему, использующую наименьшее общее число контактных пружин. Это требование иногда приводит к решению, не отвечающему предыдущему требованию, так как замыкающие и размыкающие элементы могут быть скомбинированы в переключающие элементы таким образом, что схемы, в которых они сгруппированы попарно в реле, будут удовлетворять условию 2, но не обязательно условию 1.

3. Распределить контактные пружины между всеми реле или между реле из некоторого подмножества настолько равномерно, насколько это возможно. Можно пытаться, например, найти схему, в которой нагрузка на наиболее нагруженное реле была бы по возможности минимизирована. В общем случае можно искать схему, в которой распределение нагрузки на все реле имеет некоторый специальный вид, или, настолько, насколько возможно, близко к данному распределению. Например, если реле X_1 должно срабатывать очень быстро, тогда как X_2 и X_3 не имеют существенных временных ограничений, но являются обычными реле типа U .

и X_4 — многоконтактное реле, вероятно, можно попытаться построить схему для $f(X_1, X_2, X_3, X_4)$ таким образом, чтобы прежде всего минимизировать нагрузку на X_1 , затем уравнивать нагрузку на X_2 и X_3 , делая ее в то же время настолько низкой, насколько это возможно, и, наконец, нагрузить X_4 не больше, чем это необходимо. Задачи такого типа могут быть названы задачами *распределения нагрузок*.

Хотя все эквивалентные параллельно-последовательные схемы, реализующие данную функцию f , могут быть найдены при помощи булевой алгебры, схема, отвечающая любому из указанных выше ограничений, часто может не быть параллельно-последовательной. Задача синтеза схем, не являющихся параллельно-последовательными, чрезвычайно трудна. Еще труднее показать, что схема, построенная некоторым способом, является наиболее экономичной реализацией данной функции. Сложность возникает из-за того, что имеется большое число существенно различных допустимых схем и в особенности из-за отсутствия простого математического языка для описания таких схем.

Ниже дается описание нового метода синтеза, при помощи которого может быть реализована любая функция $f(X_1, X_2, \dots, X_n)$ и часто с значительной экономией элементов по сравнению с другими методами, в частности, когда число n переменных велико. Схемы, полученные этим методом, не будут, вообще говоря, параллельно-последовательного типа, и фактически они будут даже не плоскими. Метод представляет теоретический интерес, а также пригоден и в практических целях; он дает возможность получить новые верхние оценки некоторых числовых функций, связанных с релейными схемами. Введем следующие определения:

$\lambda(n)$ определяется как наименьшее число, такое, что любая функция от n переменных может быть реализована схемой не более чем с $\lambda(n)$ элементами¹⁾. Другими словами, любая функция n переменных может быть реализована схемой не более чем с $\lambda(n)$ элементами, но найдется по крайней мере одна функция, которую нельзя реализовать схемой с меньшим числом элементов.

$\mu(n)$ определяется как наименьшее число, такое, что для любой функции f от n переменных существует двухполюсная схема, имеющая сопротивление f и использующая не больше чем $\mu(n)$ элементов в наиболее нагруженном реле.

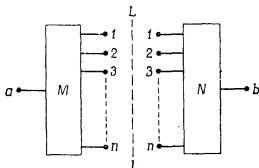
В первой части этой статьи дается общий метод синтеза схем и изучается поведение $\lambda(n)$. Во второй части изучаются допустимые распределения нагрузок на реле, а в третьей части рассматриваются

¹⁾ Элемент означает замыкающий или размыкающий контакт одного реле. Переключающий элемент означает замыкающе-размыкающий контакт с общей пружиной, содержащий два элемента.

отдельные классы функций, которые особенно легко реализуются, и доказывается несколько теорем о переключательных схемах и функциях.

2. Основная теорема синтеза

Упомянутый выше метод синтеза базируется на простой теореме, относящейся к способу соединения двух переключательных схем. Сначала сформулируем и докажем эту теорему. Допустим, что M и N (рис. 6) — две $(n + 1)$ -полюсные схемы, M реализует функцию



Р и с. 6. Схема для общей теоремы синтеза.

сопротивления U_k ($k = 1, 2, \dots, n$) между полюсами a и k и N — функцию V_k между b и k . Далее, пусть M такова, что

$$U_{jk} = 1 \quad (j, k = 1, 2, \dots, n).$$

Мы говорим в таком случае, что M есть *разделительная* схема. При этих условиях докажем следующее утверждение.

Теорема 1. Если соответствующие полюсы $1, 2, \dots, n$ схем M и N соединены между собой, то

$$U_{ab} = \prod_{k=1}^n (U_k + V_k), \quad (1)$$

где U_{ab} — сопротивление между полюсами a и b .

Доказательство. Известно, что сопротивление U_{ab} может быть найдено путем взятия произведения сопротивлений всех возможных путей от a к b , проходящих через элементы схемы¹⁾. Можно разделить эти пути на такие, которые пересекают линию L один раз, трижды, пять раз и т. д. Пусть произведение сопротивле-

¹⁾ Shannon C., A symbolic analysis of relay and switching circuits, Trans. AIEE, 57 (1938), 713. [См. наст. сборник, стр. 9. — Прим. ред.]

ний путей первого класса будет W_1 , второго — W_3 и т. д. Тогда

$$U_{ab} = W_1 \cdot W_3 \cdot W_5 \dots \quad (2)$$

Ясно, что $W_1 = \prod_{k=1}^n (U_k + V_k)$.

Кроме того,

$$W_3 = W_5 = \dots = 1,$$

так как каждый член в каждом из произведений $W_3, W_5, \dots$ должен содержать слагаемое типа U_{jk} , которое по предположению есть 1. Подставляя эти значения в (2), получим желаемый результат.

Метод использования этой теоремы при синтезе схем может в общих чертах быть описан следующим образом: реализуемая функция записывается в виде произведения типа (1) таким образом, что функции U_k одни и те же для широкого класса функций, а V_k — определяются конкретной функцией f . Основная разделительная схема M реализует функции U_k между полюсами a и k . Схема N для получения функций V_k находится специальным исследованием или по некоторым общим правилам. Рассмотрим, как это может быть выполнено в различных случаях.

3. Синтез схем для произвольных функций.

Поведение $\lambda(n)$

а. *Функции одной, двух и трех переменных.*

Функции одной или двух переменных рассматриваются легко, так как число таких функций очень мало. Так, существуют лишь следующие функции одной переменной X :

$$0, 1, X, X'$$

и, очевидно, $\lambda(1) = 1$, $\mu(1) = 1$.

От двух переменных X и Y существует 16 функций:

0	X	Y	XY	XY'	$X'Y$	$X'Y'$	$XY' + X'Y$
1	X'	Y'	$X + Y$	$X + Y'$	$X' + Y'$	$X' + Y$	$XY + X'Y'$

так что $\lambda(2) = 4$, $\mu(2) = 2$.

Пскажем, что любая функция трех переменных $f(X, Y, Z)$ может быть реализована не более чем с 8 элементами и не более чем с 4 элементами на каждом реле. Любая функция трех переменных может быть разложена в следующее произведение:

$$f(X, Y, Z) = [X + Y + f(0, 0, Z)] [X + Y' + f(0, 1, Z)] \cdot [X' + Y + f(1, 0, Z)] [X' + Y' + f(1, 1, Z)].$$

В терминах теоремы 1 полагаем

$$\begin{aligned} U_1 &= X + Y, & V_1 &= f(0, 0, Z), \\ U_2 &= X + Y', & V_2 &= f(0, 1, Z), \\ U_3 &= X' + Y, & V_3 &= f(1, 0, Z), \\ U_4 &= X' + Y', & V_4 &= f(1, 1, Z), \end{aligned}$$

так что

$$U_{ab} = f(X, Y, Z) = \prod_{k=1}^4 (U_k + V_k).$$

Приведенные выше функции U_k реализуются схемой M рис. 7, и легко видеть, что $U_{jk} = 1$ ($j, k = 1, 2, 3, 4$). Теперь нужно построить другую схему N , имеющую в качестве V_k функции V_1, V_2, V_3, V_4 . Каждая из них есть функция одной переменной Z и должна, следовательно, быть одной из четырех возможных функций одной переменной

$$0, 1, Z, Z'.$$

Рассмотрим схему N на рис. 8. Если какие-либо V равны нулю, то присоединим соответствующие полюсы схемы M к полюсу схемы N , отмеченному нулем; если какие-нибудь V равны Z , присоединим соответствующие полюсы схемы M к полюсу схемы N , отмеченному Z , и т. д. Полюсы, соответствующие 1, конечно, не соединены

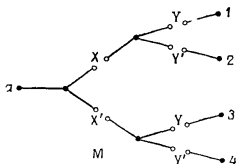


Рис. 7. Разделительное дерево с двумя ярусами.

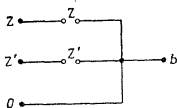


Рис. 8. Схема, представляющая функции одной переменной.

ни с чем. Из теоремы 1 следует, что схема, построенная таким образом, будет реализовать функцию $f(X, Y, Z)$. Во многих случаях некоторые из элементов будут лишними: например, если одна из V_i равна 1, элемент схемы M , соединенный с полюсом i , может быть исключен. В худшем случае M содержит 6 элементов и N содержит 2. Переменная X может встречаться дважды, Y — четырежды и Z — дважды. Конечно, совершенно несущественно, какие из переменных мы назовем X, Y или Z . Итак, доказано больше, чем было сформулировано выше, а именно следующая теорема.

Теорема 2. Любая функция трех переменных может быть реализована с использованием не более чем 2, 2 и 4 элементов соответственно для этих трех переменных, взятых в любом порядке. Таким образом, $\lambda(3) \leq 8$, $\mu(3) \leq 4$. Кроме того, так как замыкающие и размыкающие элементы содержатся в связанных парах, то в терминах переключающих элементов может быть получено распределение 1, 1, 2.

Эта теорема дает лишь верхние оценки для $\lambda(3)$ и $\mu(3)$. Возникает вопрос о том, можно ли при помощи какого-либо другого метода синтеза снизить эти оценки, т. е. можно ли знак $\leq$ заменить знаком $<$. Можно показать изучением специальных случаев, что $\lambda(3) = 8$ есть функция

$$X \oplus Y \oplus Z = X(YZ + Y'Z') + X'(YZ' + Y'Z),$$

требующая в наиболее экономичной реализации 8 элементов. $\mu(3)$, однако, фактически равно 3.

Возможно, что вообще функция

$$X_1 \oplus X_2 \oplus \dots \oplus X_n$$

требует $4(n-1)$ элементов, но это пока не доказано^{1*)}¹⁾. Доказательство того, что некоторая функция не может быть реализована с малым числом элементов, отчасти похоже на доказательство существования трансцендентных чисел. Ниже будет показано, что почти все²⁾ функции требуют для реализации большого числа элементов, однако трудно показать, что это справедливо для конкретной функции.

б. Функции четырех переменных.

При реализации функций четырех переменных тем же самым методом открываются два пути. Во-первых, можно разложить функцию следующим образом:

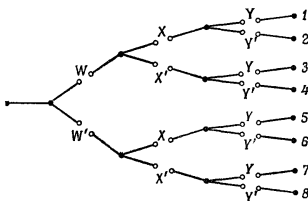
$$\begin{aligned} f(W, X, Y, Z) = & [W + X + Y + V_1(Z)] \cdot [W + X + Y' + V_2(Z)] \cdot \\ & \cdot [W + X' + Y + V_3(Z)] [W + X' + Y' + V_4(Z)] \cdot \\ & \cdot [W' + X + Y + V_5(Z)] [W' + X + Y' + V_6(Z)] \cdot \\ & \cdot [W' + X' + Y + V_7(Z)] [W' + X' + Y' + V_8(Z)]. \end{aligned}$$

При таком разложении положим $U_1 = W + X + Y$, $U_2 = W + X + Y'$, ..., $U_8 = W' + X' + Y'$ и построим схему, изображенную на рис. 9.

¹⁾ Здесь и далее цифра со звездочкой относится к примечаниям редактора в конце статьи. — *Прим. ред.*

²⁾ Выражение «почти все» использовано в арифметическом смысле, т. е. некоторое утверждение верно почти для всех функций от n переменных, если доля тех функций от n переменных, для которых это утверждение ложно, стремится к 0 при $n \rightarrow \infty$.

N здесь та же, что на рис. 8. Рассуждая так же, как в случае трех переменных, можно убедиться в том, что $\lambda(4) \leq 16$.



Р и с. 9. Разделительное дерево с тремя ярусами.

Используя несколько более сложный метод, можно, однако, снизить эту оценку. Пусть функция разложена следующим образом:

$$f(W, X, Y, Z) = [W + X + V_1(Y, Z)] \cdot [W + X' + V_2(Y, Z)] \cdot [W' + X + V_3(Y, Z)] \cdot [X' + X' + V_4(Y, Z)].$$

Можно использовать в качестве M схему типа, приведенного на рис. 7. Функции V теперь суть функции двух переменных Y и Z и могут быть любыми из следующих 16 функций:

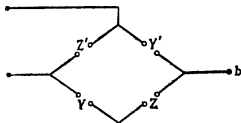
$$A \begin{Bmatrix} 0 \\ 1 \end{Bmatrix} \quad B \begin{Bmatrix} Y \\ Y' \\ Z \\ Z' \end{Bmatrix} \quad C \begin{Bmatrix} YZ \\ Y'Z \\ YZ' \\ Y'Z' \end{Bmatrix} \quad D \begin{Bmatrix} Y+Z \\ Y+Z' \\ Y'+Z \\ Y'+Z' \end{Bmatrix} \quad E \begin{Bmatrix} Y'Z + YZ' \\ YZ + Y'Z' \end{Bmatrix}.$$

Мы разделили функции на 5 групп A , B , C , D и E для дальнейших ссылок. Покажем, что любая функция четырех переменных может быть реализована не более чем с 14 элементами. Это означает, что мы должны построить схему N , использующую не более чем 8 элементов (так как в схеме M использованы 6) для любого выбора четырех функций из перечисленных выше. Для доказательства рассмотрим несколько специальных случаев.

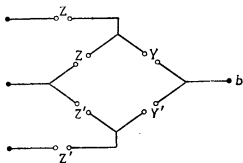
1. Если все 4 функции взяты из групп A , B , C и D , то N заведомо будет содержать не более 8 элементов, так как эти 4 функции содержат самое большее 8 символов переменных.

2. Предположим теперь, что в точности одна функция принадлежит группе E ; не нарушая общности, можно считать, что это $YZ' + Y'Z$, так как другая функция может быть получена из нее просто заменой Y на Y' . Если одна или более из оставшихся функций

принадлежат группам A или B , — ситуация удовлетворительна, для реализации этой функции не требуется дополнительных элементов. Очевидно, что для 0 и 1 элементов не нужно, а функция Y , Y' , Z или Z' может быть «извлечена» из схемы для $YZ' + Y'Z$, если последнюю функцию представить в виде $(Y+Z)(Y'+Z')$. Например, Y' можно получить при помощи схемы, изображенной на рис. 10. Оставшихся четырех элементов, несомненно, достаточно для реализации любых двух функций из групп A , B , C или D .



Р и с. 10. Упрощенная схема.



Р и с. 11. Упрощенная схема.

3. Теперь, все еще предполагая, что имеем одну функцию, $YZ' + Y'Z$, из группы E , допустим, что по крайней мере две оставшиеся функции принадлежат группе D . Используя подобный процесс «извлечения» [из схемы для $YZ' + Y'Z$], можно сэкономить по одному элементу для реализации каждой из этих функций. Например, если эти функции суть $Y + Z$ и $Y' + Z'$, то схема будет такой, как показано на рис. 11.

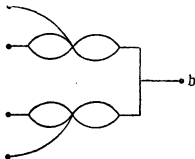
4. При том же предположении наихудшим случаем является тот, когда две функции взяты из группы C и одна из D , или все три из C .

Этот последний случай удовлетворителен, так как по крайней мере одна из этих трех функций должна быть членом функции $YZ' + Y'Z$ и может быть «извлечена». Первый из названных случаев неудовлетворителен только тогда, когда две функции из группы C суть YZ и $Y'Z'$. Легко видеть, что существенно различными являются только такие возможности для функций из группы D : $Y + Z$ и $Y' + Z$. То, что получающиеся здесь функции f могут быть реализованы с 14 контактами, может быть показано выписыванием типичных функций и упрощением их с помощью правил булевой алгебры^{2*}).

5. Рассмотрим теперь случай, где две функции принадлежат группе E . Воспользовавшись схемой рис. 12, можно «извлечь» функции или части функций из A , B или D и увидеть, что трудными случаями являются лишь следующие.

а. Две функции принадлежат группе C . В этом случае или функция f симметрична относительно Y и $Z^{3*})$ или же обе эти функции могут быть получены из схемы для функции из группы E , изображенной на рис. 12. Случай симметрии будет рассмотрен в последней части этой работы.

б. Одна функция принадлежит группе C , другая — D , причем имеет место лишь один случай отсутствия симметрии. Предположим,



Р и с. 12. Упрощенная схема.

что эти четыре функции суть $Y \oplus Z$, $Y \oplus Z'$, YZ и $Y + Z'^{4*})$. Это приводит к нескольким типам функций f , которые все могут быть упрощены с помощью алгебраических методов^{2*)}. Этим завершается доказательство теоремы, которая может быть сформулирована следующим образом:

Теорема 3. Любая функция четырех переменных может быть реализована с использованием не более чем 14 элементов^{5*)}.

в. Функции более чем четырех переменных.

Любая функция пяти переменных может быть записана в виде

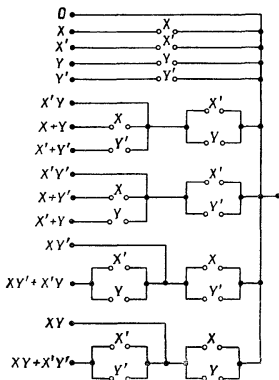
$$f(X_1, \dots, X_5) = [X_5 + f_1(X_1, \dots, X_4)] \cdot [X'_5 + f_2(X_1, \dots, X_4)], \quad (3)$$

и, поскольку, как только что было показано, каждая из двух функций от четырех переменных может быть реализована с использованием 14 элементов, $f(X_1, \dots, X_5)$ можно реализовать с 30 элементами^{6*)}.

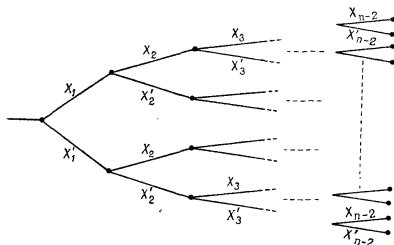
Теперь рассмотрим функцию $f(X_1, \dots, X_n)$ от n переменных. Для $5 < n \leq 10$ получим наилучшую оценку путем разложения f по всем переменным, кроме двух

$$f(X_1, \dots, X_n) = [X_1 + X_2 + \dots + X_{n-2} + V_1(X_{n-1}, X_n)] \dots \\ \dots [X'_1 + X'_2 + \dots + X'_{n-2} + V_s(X_{n-1}, X_n)]. \quad (4)$$

Функции V суть функции переменных X_{n-1} , X_n и могут быть получены из общей схемы N на рис. 13, в которой представлены все функции двух переменных.



Р и с. 13. Схема для всех функций двух переменных.



Р и с. 14. Разделительное дерево с $(n-2)$ ярусами.

Эта схема содержит 20 элементов, сгруппированных в 5 переключающих элементов для одной переменной и в 5 для другой¹⁾. Схема M для (4), изображенная на рис. 14, требует в общем случае $2^{n-1} - 2$ элементов. Таким образом, доказана теорема:

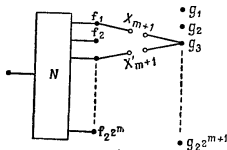
Теорема 4. $\lambda(n) \leq 2^{n-1} + 18$.

г. Верхние оценки для $\lambda(n)$ при больших значениях n .

Конечно, не слишком часто приходится реализовывать функцию более чем 10 переменных, но представляет теоретический интерес определить поведение функции $\lambda(n)$ при больших значениях n настолько точно, насколько это возможно.

Докажем сначала теорему, относящуюся к оценке числа элементов, требуемых для схемы, аналогичной схеме рис. 13, но обобщенной на случай m переменных.

Теорема 5. *Схема N , реализующая все 2^{2^m} функций m переменных^{7*)}, может быть построена с использованием менее $2 \cdot 2^{2^m}$ элементов, т. е. не более двух элементов на функцию. Любая схема с этим свойством использует по крайней мере $\left(\frac{3}{2} - \varepsilon\right)$ элементов на функцию для любого $\varepsilon > 0$ при достаточно большом значении n .*



Р и с. 15. Схема для всех функций от $(m+1)$ переменных, построенная из схемы для всех функций от m переменных.

Первая часть будет доказана по индукции. Очевидно, что это верно для $m = 1, 2$. Предположим, что это верно для некоторого m и схемы N на рис. 15. Любая функция от $m+1$ переменных может быть записана так:

$$g = [X_{m+1} + f_a] [X'_{m+1} + f_b],$$

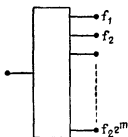
¹⁾ Можно построить несколько других схем, обладающих теми же свойствами, что и схема на рис. 13, но все они содержат те же 20 элементов.

где f_a и f_b зависят только от m переменных. Легко видеть, что все функции g могут быть получены таким же образом, как g_3 (рис. 15), т. е. путем соединения соответствующих полюсов^{8*)} f_a и f_b с полюсом g элементами X и X' . Среди них будет 2^{2^m} функций (f^{9*}), но они могут быть получены просто путем соединения с рассматриваемыми функциями f без каких-либо дополнительных элементов. Таким образом, во всей схеме используется меньше чем

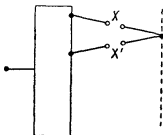
$$(2^{2^{m+1}} - 2^{2^m})2 + 2 \cdot 2^{2^m}$$

элементов, ибо по предположению схема N содержит меньше чем $2 \cdot 2^{2^m}$ элементов, и первый член в этом выражении есть число дополнительных элементов.

Второе утверждение теоремы 5 может быть доказано следующим образом. Пусть имеется схема (рис. 16), обладающая требуемым



Р и с. 16. Схема для всех функций от m переменных.



Р и с. 17. Возможный вариант в схеме на рис. 16.

свойством. Полюсы могут быть разделены на три класса: те, которые имеют один или менее элементов, непосредственно с ними соединенных; имеющие два элемента, а также те, которые имеют три или более элементов. Первое множество состоит из 0,1 и функций типа

$$(X + f) = X + f_{X=0},$$

где X есть некоторая переменная или ее отрицание. Число таких функций не превосходит $2m \cdot 2^{2^{m-1}}$ для имеющихся $2m$ способов выбора X и $2^{2^{m-1}}$ различных функций $f_{X=0}$ от оставшихся $m-1$ переменных. Следовательно, доля полюсов этого класса стремится к нулю при $n \rightarrow \infty$. Функции второго класса имеют вид

$$g = (X + f_1)(Y + f_2).$$

В случае $X \neq Y'$ функция может быть представлена в виде

$$XY + XY'g_{X=1, Y=0} + X'Yg_{X=0, Y=1} + X'Y'g_{X=0, Y=0},$$

и таких функций имеется не более чем $(2m)(2m-2)[2^{2^{m-2}}]^3$, т. е. снова бесконечно малая доля. В случае $X = Y'$ имеем положение, показанное на рис. 17, и соединение XX' не может быть использовано для другого полюса, так как такая последовательная комбинация элементов всегда разомкнута. Внутренние концы этих элементов поэтому могут быть перемещены и присоединены к полюсам, соответствующим функциям, зависящим меньше чем от n переменных, согласно равенству

$$g = (X + f_1)(X' + f_2) = (X + f_{1_{X=0}})(X' + f_{2_{X=1}}),$$

если они уже так не присоединены. Это означает, что полюсы второго класса соединены с бесконечно малой частью всех полюсов. Тогда можно относить два элемента к каждому из этих полюсов и по крайней мере $3/2$ элемента каждому из полюсов третьей группы. Так как эти две группы исчерпывают все полюсы, за исключением доли, стремящейся к нулю при $n \rightarrow \infty$, то теорема доказана.

Если при синтезе схем, реализующих функции n переменных, оборвать дерево на $(n-m)$ -м ярусе, то дерево будет содержать $2^{n-m+1} - 2$ элемента, и можно найти схему N , содержащую менее $2^{2^m} \cdot 2$ элементов, реализующую все функции оставшихся m переменных. Следовательно,

$$\lambda(n) \leq 2^{n-m+1} - 2 + 2 \cdot 2^{2^m} < 2^{n-m+1} + 2 \cdot 2^{2^m}$$

для любого целого m . Требуется найти целое число $M = M(n)$, минимизирующее эту верхнюю оценку^{10*}).

Считая, что m изменяется непрерывно, а n фиксировано, заключаем, что функция

$$f(m) = 2^{n-m+1} + 2^{2^m} \cdot 2$$

имеет только один минимум. Этот минимум должен, следовательно, находиться между m_1 и $m_1 + 1$, где

$$f(m_1) = f(m_1 + 1),$$

т. е.

$$2^{n-m_1+1} + 2^{2^{m_1}} \cdot 2 = 2^{n-m_1} + 2^{2^{m_1+1}} \cdot 2$$

или

$$2^n = 2^{m_1+1} (2^{2^{m_1+1}} - 2^{2^{m_1}}).$$

Заметим, что m_1 не может быть целым числом, так как правая часть равенства есть степень двойки и второй член (в скобках) меньше половины первого. Следовательно, для того чтобы найти целое число M , обращающее $f(M)$ в минимум, надо взять в каче-

стве M наименьшее целое число, удовлетворяющее условию^{11*})

$$2^n < 2^{M+1} \cdot 2^{2^{M+1}}.$$

Таким образом, M удовлетворяет неравенствам

$$M + 1 + 2^{M+1} > n \geq M + 2^M. \quad (5)$$

Это дает

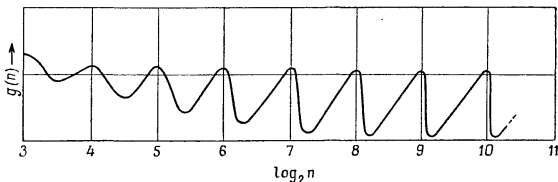
$$\begin{array}{ll} n < 11 & M = 2, \\ 11 \leq n < 20 & M = 3, \\ 20 \leq n < 37 & M = 4, \\ 37 \leq n < 70 & M = 5, \\ 70 \leq n < 135 & M = 6 \end{array}$$

и т. д.

Наша верхняя оценка для $\lambda(n)$ ведет себя подобно $2^{n+1}/n$ с пилообразными колебаниями при n , меняющимися между степенями двойки, вследствие того, что m должно быть целым числом. Если определить $g(n)$ как

$$2^{n-M+1} + 2^{2^M} \cdot 2 = g(n) \frac{2^{n+1}}{n},$$

где M определено как точка минимума функции $f(m)$ [т. е. M удовлетворяет неравенству (5)], то $g(n)$ меняется примерно так, как показано на рис. 18, где по оси абсцисс взята логарифмическая



Р и с. 18. Поведение функции $g(n)$.

шкала. Максимумы достигаются в точках, немного больших степеней двойки, причем все ближе и ближе к ним при $n \rightarrow \infty$. Пилообразный вид кривой становится все более и более точным. Резкое падение наступает при изменении одного значения M на следующее. Эти факты приводят к следующему результату.

Теорема 6^{12*} а. Для всех n

$$\lambda(n) < \frac{2^{n+3}}{n}.$$

б. Для почти всех n

$$\lambda(n) < \frac{2^{n+2}}{n}.$$

в. Существует бесконечная последовательность n_i , для которой

$$\lambda(n_i) < \frac{2^{n_i+1}}{n_i} (1 + \varepsilon), \quad \varepsilon > 0.$$

Строгое доказательство теоремы может быть получено без особых затруднений ^{13,14*}.

д. Нижняя оценка для $\lambda(n)$ при больших значениях n ,

До сих пор большая часть нашей работы относилась к определению верхней оценки для $\lambda(n)$. Было показано, что для всех n

$$\lambda(n) < B \frac{2^n}{n}.$$

Теперь спрашивается, приближаются ли где-нибудь значения этой функции $B \frac{2^n}{n}$ к значениям функции $\lambda(n)$ или, быть может, $\lambda(n)$ можно мажорировать функцией с меньшим порядком роста, например n^p . В течение некоторого времени считали, что на самом деле $\lambda(n)$ может быть мажорирована функцией n^2 для всех n , что справедливо для первых нескольких значений: 1, 4, 8, 14. Покажем, что это весьма далеко от истины и что на самом деле $\frac{2^n}{n}$ есть действительный порядок роста $\lambda(n)$:

$$A \frac{2^n}{n} < \lambda(n) < B \frac{2^n}{n}$$

для всех n . Дадим частичный ответ на следующий тесно связанный с этим вопрос, а именно: предположим, что «сложность» данной функции f от n переменных будет определена как отношение числа элементов в наиболее экономичной ее реализации к $\lambda(n)$. Тогда сложность любой функции заключена между 0 и 1. Возникает вопрос, каких функций больше — простых или сложных?

Теорема 7. Для всех достаточно больших n все функции n переменных, за исключением некоторой доли δ , требуют по крайней мере $(1 - \varepsilon) \frac{2^n}{n}$ элементов, где ε и δ — произвольно малые положительные числа. Следовательно, для больших n

$$\lambda(n) > (1 - \varepsilon) \frac{2^n}{n},$$

и почти все функции имеют сложность, большую $\frac{1}{4}(1-\varepsilon)$. Для некоторой последовательности n_i почти все функции имеют сложность, большую $\frac{1}{2}(1-\varepsilon)$.

Доказательство этой теоремы довольно интересно; оно является чистым доказательством существования. Не будем показывать, что какая-то индивидуальная функция или множество функций требуют для своей реализации $(1-\varepsilon)\frac{2^n}{n}$ элементов, но покажем, что невозможно все функции реализовать с меньшим числом элементов. Сделаем это, показав, что нельзя обойтись сетями ^{15*)} с менее чем $(1-\varepsilon)\frac{2^n}{n}$ ребрами для того, чтобы представить все 2^{2^n} функций n переменных (учитывая, конечно, различные способы приписывания переменных ребрам каждой сети). Это оказывается возможным только из-за очень быстрого роста функции 2^{2^n} . Нам потребуется следующее утверждение:

Л е м м а. Число двухполюсных сетей не более чем с K ребрами меньше, чем $(6K)^K$.

Любая двухполюсная сеть с K или меньше ребрами может быть построена следующим образом: сначала расположим по порядку K ребер и два полюса a и b :

$$\begin{array}{rcl} & 1-1' & \\ a. & 2-2' & \\ & \vdots & \\ b. & K-K' & \end{array}$$

Соединим прежде всего полюсы $a, b, 1, 2, \dots, K$ вместе любым способом ^{16*)}. Число различных способов соединения можно оценить сверху числом всех разбиений $K+2$ элементов, которое в свою очередь меньше чем 2^{K+1} , что есть число способов, которыми можно расставить один или более знаков раздела между символами $a, 1, \dots, K, b$. Далее, предполагая, что $a, 1, \dots, K, b$ соединены желаемым способом, можно соединить $1'$ или с одним из этих узлов или с дополнительным узлом, т. е. $1'$ имеет выбор самое большее из $K+3$ узлов, $2'$ имеет выбор из $K+4$ и т. д. Следовательно, число сетей заведомо меньше, чем

$$2^{K+1}(K+3)(K+4)\dots(2K+2) < (6K)^K \text{ при } K \geq 3,$$

и лемма легко проверяется для $K=1, 2$.

Вернемся теперь к доказательству теоремы 7. Число функций n переменных, которые могут быть реализованы с $\left(\frac{1-\varepsilon}{n}\right)2^n$ элементами, заведомо меньше, чем число сетей, которые можно построить

со столькими ребрами, умноженное на число способов приписывания переменных ребрам, т. е. меньше чем

$$H = (2n)^{(1-\varepsilon)\frac{2^n}{n}} \left[6(1-\varepsilon)\frac{2^n}{n} \right]^{(1-\varepsilon)\frac{2^n}{n}}.$$

Следовательно,

$$\begin{aligned} \log_2 H &= (1-\varepsilon)\frac{2^n}{n} \log_2 2n + (1-\varepsilon)\frac{2^n}{n} \log (1-\varepsilon)\frac{2^n}{n} \cdot 6 = \\ &= (1-\varepsilon)2^n + \text{члены, малые по сравнению с } (1-\varepsilon)2^n \text{ при больших } n. \end{aligned}$$

Выбирая n настолько большим, что $\frac{\varepsilon}{2} \cdot 2^n$ превосходит эти остальные члены в выражении для $\log H$, приходим к неравенству

$$\begin{aligned} \log_2 H &< (1-\varepsilon_1)2^n, \\ H &< 2^{(1-\varepsilon_1)2^n}. \end{aligned}$$

Однако имеется $S = 2^{2^n}$ функций n переменных^{7*)} и

$$\frac{H}{S} = \frac{2^{(1-\varepsilon_1)2^n}}{2^{2^n}} \rightarrow 0 \quad \text{при } n \rightarrow \infty.$$

Следовательно, почти все функции требуют больше чем $(1-\varepsilon_1)\frac{2^n}{n}$ элементов.

Далее, так как для всех $n > N$ найдется по крайней мере одна функция, требующая больше чем (скажем) $\frac{1}{2} \cdot \frac{2^n}{n}$ элементов, и так как $\lambda(n) > 0$ при $n > 0$, можно утверждать, что для всех n

$$\lambda(n) > A \frac{2^n}{n}$$

для некоторой константы $A > 0$; ибо необходимо только выбрать A таким образом, чтобы оно было наименьшим числом в конечном множестве

$$\frac{1}{2}, \frac{\lambda(1)}{2^1}, \frac{\lambda(2)}{2^2}, \dots, \frac{\lambda(n)}{2^N}.$$

Таким образом, $\lambda(n)$ есть величина порядка $\frac{2^n}{n}$. Остальные части теоремы 7 легко вытекают из уже доказанного.

По мнению автора, почти все функции имеют сложность, близкую к 1, т. е. большую $1-\varepsilon$. Это может быть показано по крайней мере для бесконечной последовательности n_i , если утверждение леммы может быть усилено — именно если будет показано, что число сетей с K ребрами для больших K меньше чем $(6K)^{K/2}$ ^{17*)}.

Хотя при подсчете числа сетей с K ребрами были использованы различные методы, все они дают результат $(6K)^K$. Интересно показать, что для больших K число сетей больше, чем

$$(6K)^{K/4}.$$

Это можно сделать обращением вышеприведенных рассуждений. Пусть $f(K)$ — число сетей с K ребрами. Так как имеется 2^{2^n} функций n переменных и каждая из них может быть реализована с $(1 + \varepsilon) \frac{2^{n+2}}{n}$ элементами (n достаточно велико), то

$$f\left((1 + \varepsilon) \frac{2^{n+2}}{n}\right) (2n)^{\frac{(1+\varepsilon) 2^{n+2}}{n}} > 2^{2^n}$$

для больших значений n . Но, как нетрудно убедиться, предположение, что $f(K) < (6K)^{K/4}$ противоречит этому неравенству. Аналогично для бесконечной последовательности K

$$f(K) > (6K)^{K/2}.$$

Так как нет никаких очевидных соображений в пользу того, что $f(K)$ связано со степенями двойки, то, по-видимому, последнее неравенство верно для всех достаточно больших K .

Теперь можно суммировать все доказанное относительно поведения $\lambda(n)$ для больших n следующим образом: $\lambda(n)$ изменяется примерно как $2^{n+1}/n$; если положить

$$\lambda(n) = A_n \frac{2^{n+1}}{n},$$

то A_n для больших значений n заключено между $(1/2 - \varepsilon)$ и $(2 + \varepsilon)$, а для некоторой бесконечной последовательности n

$$\frac{1}{2} - \varepsilon < A_n < 1 + \varepsilon.$$

Попутно было доказано, что описанный нами метод синтеза в некотором смысле не может быть существенно улучшен. Для параллельно-последовательных схем наилучшая известная оценка¹⁾ для $\lambda(n)$ есть^{18*)}

$$\lambda(n) < 3 \cdot 2^{n-1} - 2$$

¹⁾ Риордан обратил внимание на ошибку в моем рассуждении (см. стр. 28—29 данной книги. — Прим. ред.), основанном на утверждении, что эта оценка фактически достигается на функции $X_1 \oplus X_2 \oplus \dots \oplus X_n$, и показал, что эта функция и ее отрицание могут быть реализованы примерно с n^2 элементами. Ошибка содержится в разд. 4 после равенства (19) и состоит в предположении, что данное там разложение наилучшее.

и почти все функции требуют для своей реализации $(1 - \epsilon) \frac{2^n}{\log_2 n}$ элементов¹⁾. По сравнению с этим случаем имеем понижение порядка в бесконечное число раз: по крайней мере в $n/\log_2 n$ раз и, возможно, в n раз. Самое большее, что можно будет сделать в дальнейшем, — это разделить постоянный множитель на $l \leq 4$ и для некоторых n на $l \leq 2$. Возможность такого метода синтеза схем представляется, однако, маловероятной^{19*)}. Конечно, изложенное применимо только к совершенно общему методу синтеза, т. е. методу, применимому к *любой* функции. Многие специальные классы функций могут быть высоко экономично реализованы специальными методами.

Часть II

РАСПРЕДЕЛЕНИЕ НАГРУЗКИ МЕЖДУ РЕЛЕ

4. Основные положения

Рассмотрим теперь вопрос о возможно равномерном распределении нагрузок на реле или в более общем виде — о распределении нагрузок в соответствии с некоторым заданным законом. Можно думать, что такая попытка приведет к увеличению общего числа элементов в наиболее экономичной схеме. Но это не так; покажем, что во многих случаях (фактически почти для всех функций) могут быть получены очень многие распределения нагрузок, включая близкие к равномерному, при которых общее число элементов остается минимальным. Между прочим, этот результат имеет отношение к выяснению поведения $\mu(n)$, так как, комбинируя его с предыдущими теоремами, можно показать, что $\mu(n)$ имеет порядок роста $\frac{2^{n+1}}{n^2}$ при $n \rightarrow \infty$, и получить также хорошую оценку $\mu(n)$ для малых n .

Эта задача представляет скорее математический интерес, так как связана с аддитивной теорией чисел — предметом, который ранее почти не имел применений. Рассмотрим сначала несколько простых случаев. Предположим, что функция реализована деревом рис. 9. Три переменные

$$W, X, Y$$

входят соответственно

$$2, 4, 8$$

¹⁾ R i o r d a n J., S h a n n o n C., The number of two-terminal series parallel networks, *Journ. of Math. and Phys.*, 21, 2 (1942), 83. (См. наст. сборник, стр. 46—58. — Прим. ред.)

раз или в терминах переключающих элементов¹⁾

1, 2, 4.

Переменные W , X , Y могут быть переставлены любым способом без изменения функционирования схемы. Так можно переставить

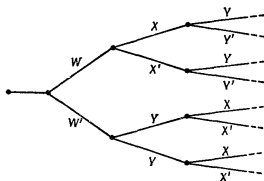


Рис. 19. Разделительное дерево с распределением 1, 3, 3.

X и Y в нижней ветви дерева, не изменяя ее функций. Получим распределение (рис. 19)

1, 3, 3.

Дерево с 4 ярусами может быть построено с любым из следующих распределений:

$W \ X \ Y \ Z$

1, 2, 4, 8 = 1, 2, 4 + 1, 2, 4,

1, 2, 5, 7 = 1, 2, 4 + 1, 3, 3,

1, 2, 6, 6 = 1, 2, 4 + 1, 4, 2,

1, 3, 3, 8 = 1, 2, 4 + 2, 1, 4,

1, 3, 4, 7 = 1, 3, 3 + 2, 1, 4,

1, 3, 5, 6 = 1, 4, 2 + 2, 1, 4,

1, 4, 4, 6 = 1, 3, 3 + 3, 1, 3,

1, 4, 5, 5 = 1, 4, 2 + 3, 1, 3

и переменные могут быть переставлены любым образом. «Суммы» справа показывают, как получены эти распределения. Первая последовательность чисел представляет верхнюю половину дерева, и вторая последовательность — нижнюю половину. Все они сводят-

¹⁾ В этой части изложение будет идти в терминах переключающих элементов.

ся к суммам последовательностей 1, 2, 4 или 1, 3, 3 в некотором порядке, а последние, как уже было отмечено, суть распределения трехъярусных деревьев. Вообще очевидно, что если можно получить распределения

$$\begin{array}{c} a_1, a_2, \dots, a_n, \\ b_1, b_2, \dots, b_n \end{array}$$

для n -ярусного дерева, то можно получить распределение

$$1, a_1 + b_1, a_2 + b_2, \dots, a_n + b_n$$

для $n + 1$ -ярусного дерева.

Заметим теперь, что все приведенные распределения обладают следующим свойством: любое из них может быть получено из первого 1, 2, 4, 8 путем перенесения одной или более единиц от больших чисел к меньшим или после последовательности таких операций, но без перенесения единицы к числу 1. Так, 1, 3, 3, 8 получается путем перенесения единицы от 4 к 2; 1, 4, 5, 5 получается путем перенесения двух единиц от 8 к 2 и одной к 4. Кроме того, каждая последовательность, которая может быть получена из последовательности 1, 2, 4, 8 посредством такого процесса, представляет собой возможное распределение. Эта операция аналогична передаче тепла — тепло может переходить только от теплого тела к холодному — так и единицы, как это показано выше, можно передавать лишь от больших чисел к меньшим.

Эти рассуждения наводят на мысль, что разделительное дерево с n ярусами может быть построено с любым распределением нагрузок, полученным при помощи такого плавного переноса единиц из начального распределения

$$1, 2, 4, 8, \dots, 2^{n-1}.$$

Покажем, что это действительно так.

Введем сначала определение. Запись $(a_1, a_2, \dots, a_n)$ обозначает любую последовательность чисел $b_1, b_2, \dots, b_n$, которая может быть получена из последовательности $a_1, a_2, \dots, a_n$ посредством следующих операций:

- 1) перестановка чисел
- 2) перенос единицы от большего числа к меньшему, не допустимый, однако, для числа 1^{20} *).

$$\begin{aligned} \text{Так,} \quad & 1, 2, 4, 8 = (1, 2, 4, 8), \\ & 4, 4, 1, 6 = (1, 2, 4, 8), \\ & 1, 3, 10, 3, 10 = (1, 2, 4, 8, 12), \end{aligned}$$

но 2, $2 \neq (1, 3)$. Условия того, что

$$b_1, b_2, \dots, b_n = (a_1, a_2, \dots, a_n), \quad (6)$$

можно представить более формально. Пусть числа a_i и b_i расположены в виде неубывающих последовательностей. Тогда необходимым и достаточным условием справедливости соотношения (6) является выполнение следующих условий:

1. $\sum_{i=1}^s b_i \geq \sum_{i=1}^s a_i \quad s = 1, 2, \dots, n;$
2. $\sum_{i=1}^n b_i = \sum_{i=1}^n a_i;$
3. существует одинаковое число единиц среди a_i и среди b_i .

Необходимость условий 2 и 3 очевидна. Условие 1 следует из того факта, что если последовательность чисел a_i неубывающая, то перенос возможен только справа налево в последовательности $a_1, a_2, \dots, a_n$

и при этом сумма $\sum_{i=1}^s a_i$ может только возрасть. Также легко установить достаточность этих условий. Если для $b_1, b_2, \dots, b_n$ выполняются условия 1, 2 и 3, то числа b_i получатся таким образом: сначала a_i увеличивается до b_1 путем последовательного перенесения единиц от чисел a_i , ближайших к a_1 (соблюдая закон «энтропии» при передаче единиц между ближайшими по величине членами), затем a_2 увеличивается до b_2 (если необходимо) и т. д. Детали достаточно очевидны^{21*}).

Аддитивная теория чисел или задача разложения числа в сумму чисел, удовлетворяющих некоторым условиям (в нашем случае это определение обобщено на «последовательности чисел»), выражается в виде следующей леммы.

Л е м м а. Если $a_1, a_2, \dots, a_n = (2, 4, 8, \dots, 2^n)$, то можно разложить последовательность чисел a_i в сумму двух последовательностей

$$a_i = b_i + c_i,$$

так что

$$b_1, b_2, \dots, b_n = (1, 2, \dots, 2^{n-1})$$

и

$$c_1, c_2, \dots, c_n = (1, 2, \dots, 2^{n-1}).$$

Можно предположить, что числа a_i образуют неубывающую последовательность $a_1 \leq a_2 \leq \dots \leq a_n$. В случае $a_1 = 2$ доказа-

тельство просто. Имеем

$$\begin{array}{ll} 1, 2, 4, \dots, 2^{n-1} & B \\ 1, 2, 4, \dots, 2^{n-1} & C \\ \hline 2, 4, 8, \dots, 2^n & A \end{array}$$

и производим перенос в последовательности $4, 8, 16, \dots, 2^n$, чтобы получить $a_2, a_3, \dots, a_n$. Любой допустимый перенос в A соответствует допустимому переносу в B или C , так как если

$$a_j = b_j + c_j > a_i = b_i + c_i,$$

то или $b_j > b_i$, или $c_j > c_i$. Поэтому для каждого переноса в сумме можно осуществить соответствующий перенос в одном или другом из слагаемых, так чтобы сумма сохранилась.

Теперь предположим, что $a_1 > 2$. Так как последовательность чисел a_i неубывающая, то

$$(n-1)a_2 \leq (2^{n+1} - 2) - a_1 \leq 2^{n+1} - 2 - 3.$$

Следовательно,

$$a_2 - 1 \leq \frac{2^{n+1} - 5}{n-1} - 1 \leq 2^{n-1}.$$

Последнее неравенство очевидно для $n \geq 5$ и легко проверяется для $n < 5$ ^{22*}). Отсюда следует, что $(a_1 - 1)$ и $(a_2 - 1)$ лежат между некоторыми степенями двойки в последовательности $1, 2, 4, \dots, 2^{n-1}$. Предположим, что

$$2^{q-1} \leq (a_1 - 1) < 2^q,$$

$$2^{p-1} \leq (a_2 - 1) < 2^p, \quad 2 \leq q \leq p \leq (n-1).$$

Будем производить перенос между 2^q и 2^{q-1} до тех пор, пока одно из них не достигнет значения $(a_1 - 1)$, а другое, скажем, R ; аналогичным образом поступим для получения $(a_2 - 1)$; другое число при этом пусть равно S . Тогда в начале нашего разложения имеем последовательности (после выполнения перестановок)

$$\begin{array}{cc|l} (a_1 - 1) & 1 & L \\ 1 & a_2 - 1 & \\ \hline a_1 & a_2 & \end{array} \left| \begin{array}{l} 2, 4, \dots, 2^{q-2}, R, \quad 2^{q+1}, \dots, 2^{p-1}, 2^p, 2^{p+1}, \dots, 2^{n-1} \\ 2, 4, \dots, 2^{q-2}, 2^{q-1}, 2^q, \dots, 2^{p-2}, S, 2^{p+1}, \dots, 2^{n-1} \\ \hline 4, 8, \dots, 2^{q-1}, \dots \qquad \qquad \qquad \dots, 2^{p+2}, \dots, 2^n \end{array} \right.$$

Теперь требуется преобразовать величины справа от $L-L$ в величины $a_3, a_4, \dots, a_n$. Обозначим последовательность

$$4, 8, \dots, 2^{q-1}, (2^{q-1} + R), 3 \cdot 2^q, 3 \cdot 2^{q+1}, \dots, (2^p + S), 2^{p+2}, \dots, 2^n$$

через $\mu_1, \mu_2, \dots, \mu_{n-2}$. Поскольку обе строки в упомянутом выше сложении последовательностей неубывающие справа от $L-L$ и не содержат единиц, то лемма будет доказана, если покажем, что

$$\sum_{j=1}^i \mu_j \leq \sum_{j=3}^{i+2} a_j \quad i = 1, 2, \dots, n-2,$$

так как уже было показано, что это достаточное условие того, что

$$a_3, a_4, \dots, a_n = (\mu_1, \mu_2, \dots, \mu_{n-2}).$$

Этим обстоятельством в первую очередь и воспользуемся. Для $i \leq q-2$, т. е. перед членом $(2^{q-1} + R)$

$$\sum_{j=1}^i \mu_j = 4(2^i - 1)$$

и

$$\sum_{j=3}^{i+2} a_j \geq ia_2 \geq i2^{p-1} \geq i2^{q-1},$$

так как

$$q \leq p.$$

Следовательно, ^{23*)}

$$\sum_{j=1}^i \mu_j \leq \sum_{j=3}^{i+2} a_j, \quad i \leq q-2.$$

Далее, для $(q-1) \leq i \leq (p-2)$, т. е. перед членом $(2^p + S)$

$$\sum_{j=1}^i \mu_j = 4(2^{q-2} - 1) + 2^{q-1} + R + 3 \cdot 2^q(2^{i-q+1} - 1) < 3 \cdot 2^{i+1} - 4 \leq 3 \cdot 2^{p-1} - 4.$$

Так как $R < 2^q$, то снова имеем

$$\sum_{j=3}^{i+2} a_j \geq i2^{p-1},$$

так что в этом интервале также имеем желаемое неравенство ^{24*)}. Наконец, для последнего интервала

$$\sum_{j=1}^i \mu_j = 2^{i+2} - a_1 - a_2 \leq 2^{i+3} - a_1 - a_2 - 2$$

и

$$\sum_{j=3}^{i+2} a_j = \sum_{j=1}^{i+2} a_j - a_1 - a_2 \geq 2^{i+3} - a_1 - a_2 - 2,$$

так как

$$a_1, a_2, \dots, a_n = (2, 4, 8, \dots, 2^n).$$

Тем самым лемма доказана.

5. Разделительное дерево

Теперь легко доказать следующее.

Теорема 8. *Разделительное дерево с n ярусами может быть построено с любым распределением*

$$a_1, a_2, \dots, a_n = (1, 2, 4, \dots, 2^{n-1}).$$

Можно доказать это утверждение по индукции. Как было показано, оно верно для $n=2, 3, 4$. Из предположения, что оно верно для n , следует, что оно верно для $n+1$, поскольку лемма утверждает, что любая последовательность

$$a_1, a_2, \dots, a_n = (2, 4, 8, \dots, 2^n)$$

может быть разложена в сумму двух последовательностей, каждая из которых может быть реализована деревом в силу индуктивного предположения.

Ясно, что среди возможных распределений

$$(1, 2, 4, \dots, 2^{n-1})$$

для дерева может быть найдено «почти равномерное» распределение для всех переменных, кроме одного. Таким образом, можно распределить нагрузки между $(n-1)$ из них равномерно, а для одного переменного использовать один (переключающий) элемент. Получим такие близкие к равномерному распределения:

$n=1$	1
$n=2$	1, 2
$n=3$	1, 3, 3
$n=4$	1, 4, 5, 5
$n=5$	1, 7, 7, 8, 8
$n=6$	1, 12, 12, 12, 13, 13
$n=7$	1, 21, 21, 21, 21, 21, 21

и. т. д.

6. Другие задачи распределения

Рассмотрим теперь задачу распределения нагрузок в параллельно-последовательных схемах. Докажем следующую теорему.

Теорема 9. *Любая функция $f(X_1, X_2, \dots, X_n)$ может быть реализована параллельно-последовательной схемой с распределением*

$$(1, 2, 4, \dots, 2^{n-2}), 2^{n-2}$$

в терминах переключающих элементов.

Докажем это утверждение по индукции. Оно верно для $n = 3$, так как любая функция трех переменных может быть представлена следующим образом:

$$f(X, Y, Z) = [X + f_1(Y, Z)][X' + f_2(Y, Z)]$$

и каждая из функций $f_1(Y, Z)$ и $f_2(Y, Z)$ может быть реализована с одним переключающим элементом переменной Y и одним — переменной Z . Таким образом, $f(X, Y, Z)$ может быть реализована с распределением 1, 2, 2. Теперь, предполагая, что теорема верна для $n - 1$, имеем

$$f(X_1, X_2, \dots, X_n) = [X_n + f_1(X_1, X_2, \dots, X_{n-1})] \cdot [X'_n + f_2(X_1, X_2, \dots, X_{n-1})]$$

и

$$\begin{array}{c} 2, 4, 8, \dots, 2^{n-3} \\ 2, 4, 8, \dots, 2^{n-3} \\ \hline 4, 8, 16, \dots, 2^{n-2} \end{array}$$

Простое применение леммы дает нужный результат.

Возможны также многие другие распределения, кроме устанавливаемых теоремой 9, но еще не найдены простые критерии для их описания. Ничего нельзя сказать о любом распределении

$$(1, 2, 4, 8, \dots, 2^{n-2}, 2^{n-2})$$

(во всяком случае, на основании нашего анализа), так как, например, последовательность

$$3, 6, 6, 7 = (2, 4, 8, 8)$$

не может быть разложена в две последовательности

$$a_1, a_2, a_3, a_4 = (1, 2, 4, 4)$$

и

$$b_1, b_2, b_3, b_4 = (1, 2, 4, 4).$$

Однако, по-видимому, допустимо почти равномерное распределение.

В качестве последнего примера распределения нагрузок рассмотрим случай схемы, состоящей из нескольких деревьев от одних и тех же переменных. Большое число таких случаев будет рассмотрено позднее. Справедливость следующего утверждения достаточно очевидна из того, что уже доказано.

Теорема 10. *Можно построить m различных деревьев от одних и тех же n переменных со следующим распределением:*

$$a_1, a_2, \dots, a_n = (m, 2m, 4m, \dots, 2^{n-1}m).$$

Интересно отметить, что при этих условиях исключается случай единицы при $m > 1$. Можно уравнивать нагрузку на всех n переменных, а не только на $n - 1$ из них.

7. Функция $\mu(n)$

Теперь рассмотрим поведение функции $\mu(n)$. Это можно сделать вместе с изучением возможных распределений нагрузок в общей функции n переменных. Уже было показано, что любая функция трех переменных может быть реализована с распределением в терминах переключающих 1, 1, 2 элементов, следовательно, $\mu(3) \leq 4$.

Любая функция четырех переменных может быть реализована с распределением

$$1, 1, (2, 4).$$

Следовательно, $\mu(4) \leq 6$. Для пяти переменных можно получить распределение

$$1, 1, (2, 4, 8)$$

или

$$1, 5, 5, (2, 4),$$

так что $\mu(5) \leq 10$. Для шести переменных получаем

$$1, 5, 5, (2, 4, 8) \text{ и } \mu(6) \leq 10;$$

для семи

$$1, 5, 5, (2, 4, 8, 16) \text{ и } \mu(7) \leq 16$$

и т. д. Таким образом, поскольку возможно равномерно распределить нагрузку между всеми реле в дереве, за исключением одного, то можно сформулировать теорему, аналогичную теореме 7, для функции $\mu(n)$.

Теорема 11. Для всех n справедливо неравенство

$$\mu(n) \leq \frac{2^{n+3}}{n^2}.$$

Для почти всех n справедливо неравенство

$$\mu(n) \leq \frac{2^{n+2}}{n^2}.$$

Для бесконечного числа n_i справедливо неравенство

$$\mu(n_i) \leq (1 + \epsilon) \frac{2^{n+1}}{n^2}.$$

Доказательство прямое и здесь опускается ^{25*}).

Часть III

СПЕЦИАЛЬНЫЕ ФУНКЦИИ

8. Функциональные отношения

Было показано, что почти все функции требуют для своей реализации порядка $2^{n+1}/n^2$ элементов на реле. Тем не менее для практических целей это число слишком велико. В телеграфном аппарате, например, где реализуется много функций, некоторые из которых зависят от большого числа переменных, реле в среднем содержат 7 или 8 контактов. Фактически почти все реле, с которыми мы сталкиваемся на практике, имеют меньше 20 элементов. В чем причина этого парадокса? Ответ, конечно, состоит в том, что функции, с которыми сталкиваются на практике, далеки от случайных. Здесь снова имеется аналогия с трансцендентными числами. Хотя почти все числа трансцендентны, вероятность встретить трансцендентное число в открытой наудачу математической книге, конечно, гораздо меньше единицы. Встречающиеся обычно функции гораздо проще, чем основная масса булевых функций в силу по крайней мере двух основных причин.

1. Инженер, синтезирующий схему, имеет значительную свободу выбора функций, которые подлежат реализации, и часто может выбрать достаточно простые. Например, синтезируя передающие устройства для работы телефона, используют обычно аддитивные коды, а также коды, в которых одно и то же количество реле работает для каждой возможной цифры. Логическая простота этих кодов отражена в простоте схем, оперирующих с этими кодами.

2. Большинство условий, предъявляемых переключательным схемам, имеет логически простую природу. Наиболее важным аспектом этой простоты является то, что большинство схем может быть разбито на большое число малых схем. Вместо реализации функции от большого числа переменных реализуется много функций от малого числа переменных и затем некоторая функция от этих функций. Чтобы показать эффективность этого метода, рассмотрим следующий пример. Предположим, что требуется реализовать функцию

$$f(X_1, X_2, \dots, X_{2n})$$

от $2n$ переменных. Лучшая оценка, которую можно дать для общего числа необходимых элементов, есть приблизительно $2^{2n+1}/2n$. Однако если известно, что f есть функция от двух функций f_1 и f_2 , каждая из которых зависит только от n переменных, т. е. если

$$f = g(f_1, f_2), \quad f_1 = f_1(X_1, \dots, X_n), \quad f_2 = f_2(X_{n+1}, \dots, X_{2n}),$$

то можно реализовать f примерно с $4 \cdot 2^{n+1}/n$ элементами, что имеет

значительно меньший порядок роста, чем $2^{2^{n+1}}/2n$. Когда g есть одна из простейших функций двух переменных, например, когда $g(f_1, f_2) = f_1 + f_2$, а также в любом другом случае, когда g реализуется двумя элементами, можно реализовать f примерно $2 \cdot 2^{n+1}/n$ элементами. Вообще, чем дальше было бы возможно разложить задачу синтеза в комбинацию простых задач, тем проще окажется окончательная схема. Важным пунктом здесь является тот факт, что f удовлетворяет определенному функциональному отношению

$$f = g(f_1, f_2)$$

и можно найти более простую схему по сравнению со средней функцией того же числа переменных.

Этот тип функциональных соотношений может быть назван функциональной делимостью. Ее часто можно обнаружить из условий функционирования схемы и можно использовать для уменьшения числа требуемых элементов. Покажем теперь, что большинство функций не являются функционально делимыми.

Теорема 12. Доля функций n переменных, которые могут быть записаны в виде ^{26*}

$$f = g(h(X_1, \dots, X_s), X_{s+1}, \dots, X_n),$$

где $1 < s < n-1$, стремится к 0 при $n \rightarrow \infty$.

Можно выбрать s переменных, входящих в h , C_n^s способами; функция h при этом имеет 2^{2^s} возможностей и g имеет $2^{2^{n-s+1}}$ возможностей, поскольку зависит от $n-s+1$ аргументов. Общее число функционально делимых функций, следовательно, не превосходит

$$\sum_{s=2}^{n-2} C_n^s 2^{2^s} \cdot 2^{2^{n-s+1}} \leq (n-3) \frac{n^2}{2} 2^{2^2} \cdot 2^{2^{n-1}}$$

и отношение этого числа к 2^{2^n} стремится к 0 при $n \rightarrow \infty$.

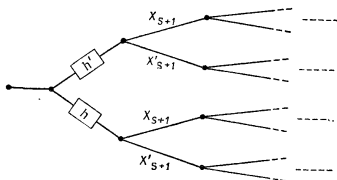
Если имеет место такая функциональная делимость, то во многих случаях с успехом может быть использован описанный выше общий метод синтеза. В общем виде это показано на рис. 20. Если делимость более сильная, например,

$$f = g[h_1(X_1, \dots, X_s), h_2(X_{s+1}, \dots, X_t), X_{t+1}, \dots, X_n],$$

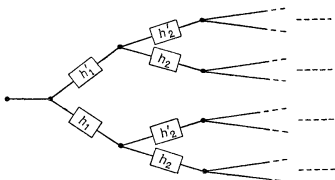
можно воспользоваться схемой рис. 21, беря в качестве h_2 ту из функций h_1 и h_2 , которая вместе со своим отрицанием требует для реализации меньшего числа элементов.

Рассмотрим теперь второй тип функциональных соотношений, часто встречающийся на практике и помогающий при выборе экономической реализации. Этот тип соотношений может быть назван груп-

повой инвариантностью; ее специальный случай — симметрия функций относительно всех переменных — был рассмотрен в цитированной на стр. 60 работе автора. Функцию $f(X_1, X_2, \dots, X_n)$ будем



Р и с. 20. Использование разделимости для уменьшения числа элементов.



Р и с. 21. Использование разделимости для двух множеств переменных.

называть симметричной относительно X_1 и X_2 , если она удовлетворяет соотношению

$$f(X_1, X_2, \dots, X_n) = f(X_2, X_1, \dots, X_n);$$

функция симметрична относительно X_1 и X'_2 , если выполнено равенство

$$f(X_1, X_2, \dots, X_n) = f(X'_2, X'_1, \dots, X_n).$$

Это частные случаи рассматриваемого типа функциональных соотношений. Обозначим через

$N_{0,0}, \dots, 0 = I$ — операцию оставления переменных в функции такими, как они есть,

$N_{1,0}, \dots, 0$ — операцию отрицания первой (т. е. стоящей на первом месте) переменной,

$N_{0,1,0,\dots,0}$ — операцию отрицания второй переменной,

$N_{1,1,\dots,0}$ — операцию отрицания первых двух переменных и т. д., так что $N_{101}f(X, Y, Z) = f(X', Y, Z')$ и т. д.

Совокупность всех N_i образует абелеву группу с тем важным свойством, что каждый элемент является обратным самому себе, $N_i N_i = I$. Произведение двух элементов можно легко найти, но $N_i N_j = N_k$, где k есть число, находящееся посредством сложения i и j как двоичных чисел без переноса.

Заметим, что эта «отрицающая» группа состоит из 2^n элементов. Пусть теперь

$S_{1,2,3,\dots,n} = I$ — операция сохранения переменных функции на прежних местах,

$S_{2,1,3,\dots,n}$ — операция перестановки первых двух переменных,

$S_{3,2,1,4,\dots,n}$ — операция перестановки первого и третьего переменных и т. д.

Так,

$$S_{3,1,2}f(X, Y, Z) = f(Z, X, Y)$$

$$S_{3,1,2}f(Z, X, Y) = S_{3,1,2}^2 f(X, Y, Z) = \text{и т. д.}$$

Совокупность всех S_i также образует группу, знаменитую группу «подстановок» или «симметрическую» группу. Она имеет порядок $n!$. Правда, она не обладает простыми свойствами отрицающей группы — она не абелева ($n > 2$) и не каждый ее элемент является обратным себе¹⁾. Отрицающая группа не является циклической, если $n \geq 2$; симметрическая группа не является циклической, если $n \geq 3$.

Полупрямое произведение этих двух групп образует группу G , общий элемент которой имеет вид $N_i S_j$, и так как i может принимать 2^n значений, а j может принимать $n!$ значений, порядок группы G равен $2^n \cdot n!$

Легко видеть, что $S_j N_i = N_k S_j$, где k получается из преобразования i , рассматриваемого как упорядоченная последовательность нулей и единиц, перестановкой S_j . Так

$$S_{2,3,1,4} N_{1,1,0,0} = N_{1,0,1,0} S_{2,3,1,4}.$$

При помощи этого правила любые произведения вида $N_i S_j N_k N_l S_m N_n S_p$ могут быть приведены к виду $N_i N_j \dots N_n S_p S_q \dots S_r$, что преобразуется к стандартному виду $N_i S_j$.

¹⁾ Второе условие лишнее; свойство каждого элемента быть себе обратным влечет коммутативность, ибо если для каждого элемента X выполнено условие $XX = I$, то $XY = (XY)^{-1} = Y^{-1}X^{-1} = YX$.

$$N_i S_i f = f.$$

Элемент группы, оставляющий функцию f инвариантной, определяет некоторые равенства для членов, входящих в разложение f . Чтобы показать это, рассмотрим фиксированный элемент $N_i S_j$, который преобразует некоторым образом переменные $X_1, X_2, \dots, X_r$. Пусть функция $f(X_1, X_2, \dots, X_n)$ разложена по переменным $X_1, X_2, \dots, X_r$:

$$\begin{aligned} f = & [X_1 + X_2 + \dots + X_r + f_1(X_{r+1}, \dots, X_n)] \\ & [X'_1 + X_2 + \dots + X_r + f_2(X_{r+1}, \dots, X_n)] \\ & \cdot \\ & [X'_1 + X'_2 + \dots + X'_r + f_{2^r}(X_{r+1}, \dots, X_n)]. \end{aligned}$$

Если f удовлетворяет соотношению $N_i S_j f = f$, то покажем, что имеется по крайней мере $\frac{1}{4} 2^r$ равенств между функциями $f_1, f_2, \dots, f_{2^r}$. Таким образом, число функций, удовлетворяющих этому соотношению, не превосходит $(2^{2^{n-r}})^{3/4 2^r} = 2^{3/4 2^n}$, поскольку каждая независимая функция f_i может быть любой точно из $2^{2^{n-r}}$ функций, и имеется самое большее $\frac{3}{4} 2^r$ независимых функций f_i . Пусть $N_i S_j$ преобразует

$$X_1, X_2, \dots, X_r \quad (A)$$

$$X_{a_1}^*, X_{a_2}^*, \dots, X_{a_r}^*, \quad (B)$$

где символ * может означать или штрих, или его отсутствие, но не может быть $X_{a_i}^* = X_i$. Дадим переменной X_i значение 0. Это фиксирует некоторый элемент в B , а именно X_{a_i} , где $a_i = 1$. Возможны два случая.

1). Если этот элемент есть первый член, $a_1 = 1$, то имеем

$$0, X_2, \dots, X_r$$

$$1, X_{a_0}, \dots, X_{a_r}.$$

Давая переменным $X_2, \dots, X_r$ все 2^{r-1} их возможных значений, полу-

чим 2^{r-1} равенств между различными функциями f_i , так как они на самом деле суть

$$f(X_1, X_2, \dots, X_r, X_{r+1}, \dots, X_n)$$

с фиксированными $X_1, X_2, \dots, X_r$.

2) Если элемент, о котором идет речь, есть другой член, скажем $X_{a_2}^*$, дадим переменной X_2 в ряде A противоположное значение $X_2 = (X_{a_2}^*)' = (X_2^*)'$. Теперь, поступая так же, как и выше, с оставшимися $r-2$ переменными, установим 2^{r-2} равенств между f_i .

Далее, имеется не более $^1)$ чем $2^n n!$ соотношений

$$N_i S_j f = f,$$

которым функция может удовлетворять, так что число функций, удовлетворяющих хотя бы одному нетривиальному соотношению, не превосходит $2^n n! 2^{3/4 2^n}$. Поскольку

$$2^n n! 2^{3/4 2^n} / 2^{2^n} \rightarrow 0 \text{ при } n \rightarrow \infty,$$

имеем следующую теорему.

Теорема 13. *Почти все функции не имеют нетривиальной группы инвариантности.*

Из теорем 12 и 13, а также из других результатов вытекает, что почти все функции имеют крайне хаотическую природу и не проявляют никакой симметрии или каких-либо других функциональных соотношений. Это можно было бы предсказать на основании того, что такие соотношения вообще ведут к значительному уменьшению числа требуемых элементов, а было показано, что почти все функции имеют довольно большую «сложность».

Если синтезировать функцию методом разделительного дерева и функция имеет группу инвариантности, включающую переменные

$$X_1, X_2, \dots, X_r,$$

то по крайней мере 2^{r-2} полюсов в соответствующем дереве могут быть соединены с другими, так как существует по крайней мере столько же равенств между функциями, соответствующими этим полюсам. Это в общем приводит к значительному уменьшению числа требуемых контактов оставшихся переменных. Можно также достигнуть некоторой экономии в схеме M . Для применения этого метода синтеза, однако, существенно то, что у нас есть способ определения, какой из элементов $N_i S_j$, если он существует, оставляет функцию неизменной. Следующая теорема, хотя это и не все,

¹⁾ Наш множитель на самом деле меньше потому, что, во-первых, надо исключить $N_i S_j = I$ и, во-вторых, за исключением обратных себе элементов, одно отношение такого типа влечет другие, а именно его степени:

$$(N_i S_j)^p f = f.$$

что можно ожидать, показывает, что вовсе не обязательно рассматривать $N_i S_j f$ для всех $N_i S_j$, а достаточно рассмотреть $N_i f$ и $S_j f$.

Теорема 14. *Необходимым и достаточным условием того, что $N_i S_j f = f$, является выполнение равенства $N_i f = S_j f$.*

Это непосредственно вытекает из свойств N_i быть обратным себе. Конечно, групповую инвариантность часто можно определить, исходя непосредственно из требований к схеме в задаче синтеза.

Таблица I

Групповая инвариантность для двух переменных (рис. 22)

	$S_{1,2}$	$S_{2,1}$
$N_{0,0}$	(XY)	$(YX)^{1*}$
$N_{0,1}$	$(XY')^{2*}$	$(YX')^{3*}$
$N_{1,0}$	$(X'Y)^2$	$(Y'X)^{3*}$
$N_{1,1}$	$(X'Y')^4$	$(Y'X')^1$

Таблица II

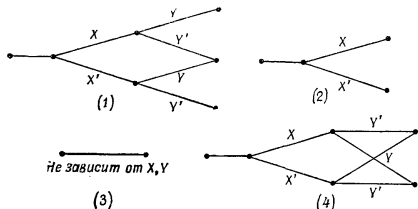
Групповая инвариантность для трех переменных (рис. 23)

	$S_{1,2,3}$	$S_{1,3,2}$	$S_{2,1,3}$	$S_{2,3,1}$	$S_{3,1,2}$	$S_{3,2,1}$
$N_{0,0,0}$	XYZ	XZY^1	YXZ^1*	YZX^{2*}	ZXY^{2*}	ZYX
$N_{0,0,1}$	XYZ'^{3*}	XZY'^{4*}	YXZ'^7	YZX'^9	ZXY'^9	ZYX'^4
$N_{0,1,0}$	$XY'Z^3$	$XZ'Y^4*$	$YX'Z^4$	$YZ'X^9$	$ZX'Y^9$	$ZY'X^7$
$N_{0,1,1}$	$XY'Z'^5$	$XZ'Y'^1$	$YX'Z'^8$	$YZ'X'^2$	$ZX'Y'^2$	$ZY'X'^8$
$N_{1,0,0}$	$X'YZ^3$	$X'ZY^7*$	$Y'XZ^4$	$Y'ZX^9$	$Z'XY^9$	$Z'YX^4$
$N_{1,0,1}$	$X'YZ'^5$	$X'ZY'^{8*}$	$Y'XZ'^{8*}$	$Y'ZX'^2$	$Z'XY'^2$	$Z'YX'$
$N_{1,1,0}$	$X'Y'Z^5*$	$X'Z'Y^8*$	$Y'X'Z^1$	$Y'Z'X^2$	$Z'X'Y^2$	$Z'Y'X^{8*}$
$N_{1,1,1}$	$X'Y'Z'^6$	$X'Z'Y'^7$	$Y'X'Z'^7$	$Y'Z'X'^9$	$Z'X'Y'^9$	$Z'Y'X'^7$

Таблицы I и II построены для тех случаев, когда инвариантность имеет место для двух или трех переменных. Чтобы показать, как пользоваться этими таблицами, предположим, что задана функция, такая, что

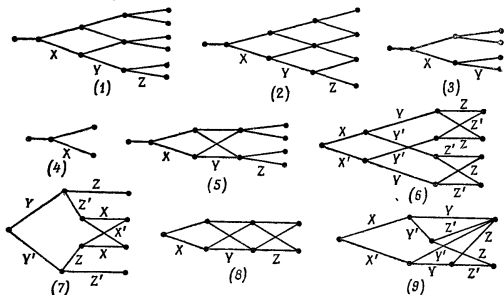
$$N_{1,1,1} S_{3,1,2} f = f.$$

Соответствующее место $Z'X'Y'$ групповой таблицы отсылает нас к схеме 9 рис. 23. Звездочки показывают, что схема может быть



Р и с. 22. Схемы групповой инвариантности двух переменных.

использована непосредственно; отсутствие их означает, что требуется перестановка переменных. Разложим f по переменным X, Y, Z ; только две различные функции будут содержаться в сомножителях. Эти две функции реализуются двумя деревьями, исходящими



Р и с. 23. Схемы групповой инвариантности трех переменных.

из полюсов схемы 9. Любая такая функция f может быть реализована схемой (с использованием одного переменного в схеме N) с

$$10 + 2(2^{n-3} - 2) + 2 = 2^{n-2} + 8 \text{ элементами;}$$

оценка гораздо лучше, чем $2^{n-1} + 18$ для произвольной функции.

9. Частично симметрические функции

Будем говорить, что функция является «частично симметрической» или «симметрической относительно некоторого множества переменных», если эти переменные могут быть переставлены без изменения функции. Так, функция

$$XYZW + (XY' + X'Y)W + WZ'$$

симметрическая относительно X и Y . Это свойство, очевидно, представляет собой частный случай общей групповой инвариантности, которая уже рассмотрена. Известно, что любая функция, симметрическая относительно всех переменных, может быть реализована схемой не более чем с n^2 элементами, где n есть число переменных¹⁾. В данном разделе будет усилен и обобщен этот результат.

Теорема 15. Любая функция $f(X_1, X_2, \dots, X_m, Y_1, \dots, Y_n)$, симметрическая относительно $X_1, X_2, \dots, X_m$, может быть записана в виде

$$\begin{aligned} f(X_1, X_2, \dots, X_m, Y_1, Y_2, \dots, Y_n) = & \\ = [S_0(X_1, X_2, \dots, X_m) + f_0(Y_1, Y_2, \dots, Y_n)] \times & \\ \times [S_1(X_1, X_2, \dots, X_m) + f_1(Y_1, Y_2, \dots, Y_n)] \times & \\ \dots \dots \dots & \\ \times [S_m(X_1, X_2, \dots, X_m) + f(Y_1, Y_2, \dots, Y_n)], & \end{aligned} \quad (6)$$

где $f_k(Y_1, Y_2, \dots, Y_n) = f(\underbrace{0, 0, \dots, 0}_k, \underbrace{1, 1, \dots, 1}_{(n-k) \text{ единиц}}, Y_1, Y_2, \dots, Y_n)$

и $S_k(X_1, X_2, \dots, X_m)$ — симметрическая функция от $X_1, X_2, \dots, X_m$ с a -числом, равным k .

Эта теорема следует из того факта, что поскольку f — симметрическая функция относительно $X_1, X_2, \dots, X_m$, то значение функции f зависит только от числа X_i , которые равны нулю, и значений Y . Если в точности k X_i равны 0, то значение f есть f_k , но правая часть равенства (6) в этом случае также обращается в f_k , так как тогда $S_j(X_1, X_2, \dots, X_m) = 1$, $j \neq k$ и $S_k = 0$.

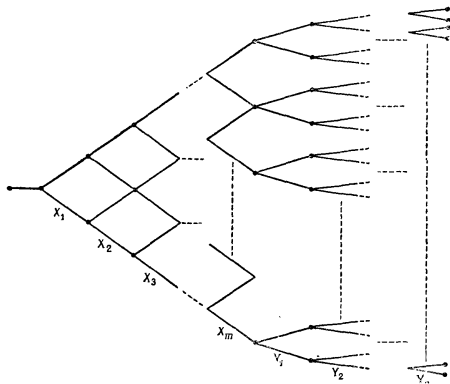
Разложение (6) имеет вид, удобный для нашего метода синтеза. Можно реализовать разделительные функции $S_k(X_1, X_2, \dots, X_m)$ симметрической решеткой и продолжить обычными деревьями (рис. 24), по одному дереву для каждого уровня схемы для симметрических функций. Обрывая деревья на ярусе с переменной Y_{n-1} , видим, что полная схема разделительная, и второе применение теоремы 1 позволяет нам реализовать функцию с двумя элементами для Y_n . Таким образом, имеем следующую теорему.

¹⁾ См. работу на стр. 9 данного сборника.

Теорема 16. Любая функция $m + n$ переменных, симметрическая относительно m из них, может быть реализована с числом элементов, не превосходящим меньшее из двух чисел

$$(m + 1)(\lambda(n) + m) \text{ или } (m + 1)(2^n + m - 2) + 2.$$

В частности, функция n переменных, симметрическая относительно $n - 2$ или более из них, может быть реализована не более чем с $n^2 - n + 2$ элементами.



Р и с. 24. Схема M для частично симметрических функций $N_{1,1,1}S_{3,1,2}f = f$.

Если функция симметрическая относительно $X_1, X_2, \dots, X_m$, а также относительно $Y_1, Y_2, \dots, Y_r$, но не симметрическая относительно $Z_1, Z_2, \dots, Z_n$, то она может быть реализована тем же самым методом с использованием симметрической решетки вместо дерева для переменных Y . Эта функция сначала разлагается по X (в предположении, что $m < r$), затем по Y и, наконец, по Z . Часть схемы, содержащая переменные Z , будет состоять из $(m + 1)(r + 1)$ деревьев^{27*)}.

ПРИМЕЧАНИЯ РЕДАКТОРА

1*) Это утверждение впоследствии было доказано Кардо. (Cardot C., Quelques résultats sur l'application de l'algèbre de Boole à la synthèse des circuits à relais, *Ann. Telecommunic.*, 7, № 1—2 (1952) 75.)

2*) Достаточно, например, разложить получающиеся функции по переменным Y, Z .

3*) Здесь имеется в виду симметрия при $Y \leftrightarrow Z$ или при $Y \leftrightarrow Z'$.

4*) Можно считать, что функция из C есть именно YZ . Если при этом функция из D есть $Y+Z$ или $Y'+Z'$, то имеет место симметрия относительно Y и Z .

5*) Г. Н. Поваров составил каталог контактных схем для всех 402 типов функций 4 переменных (Поваров Г. Н., Исследование контактных схем с минимальным числом контактов. Диссертация, Москва, 1954). Другой каталог (для типов функций 4 переменных, обращающихся в 0 на 8 и менее наборах) составили Р. Игонне и Р. Греа. (Higonnet R., Gréa R., Etude logique des circuits électriques et des systèmes binaires, Paris, 1955.) Ю. Л. Васильев, исходя из этих каталогов и упростив более 100 схем, получил минимальные схемы для всех 402 типов функций 4 переменных (Васильев Ю. Л., Минимальные контактные схемы для булевых функций четырех переменных, *ДАН*, 1959, 127, № 2, 242). При этом, в частности, Васильевым установлено, что $\lambda(4) = 13$.

6*) Пользуясь своим каталогом, Г. Н. Поваров показал, что $\lambda(5) \leq 28$.

7*) Точнее, переменных $X_1, X_2, \dots, X_n$.

8*) Здесь и ниже в доказательстве теоремы отождествляются функции и полюсы, на которых они реализуются.

9*) От m переменных.

10*) Для получения асимптотического неравенства $\lambda(n) \leq \frac{2^{n+2}}{n}$ достаточно в качестве $M(n)$ взять $\lceil \log_2(n-2 \log_2 n) \rceil$.

11*) Пусть $f_n(m) = 2^{2^m} + 2^{n-m}$. Тогда если

$$M + 2^M \leq n < M + 1 + 2^{M+1}, \quad (*)$$

то

$$f_n(M-1) > f_n(M) \leq f_n(M+1).$$

В самом деле, из левой части неравенства (*) следует, что $2^M \leq n - M$. Поэтому $f_n(M) = 2^{2^M} + 2^{n-M} \leq 2^{2^M} + 2^{n-M+1}$; $f_n(M-1) = 2^{2^{M-1}} + 2^{n-M+1} \geq 2^{2^M} + 2^{n-M+1}$. Наконец, если $n = M + 2^{M+1}$, то $f_n(M) = 2^{2^M} + 2^{2^M+1}$ и $f_n(M+1) = 2^{2^{M+1}} + 2^{2^M+1} - 1 \geq f_n(M)$; так как $2^{2^M+1} - 1 \geq 2^{2^M}$. Если же $n \leq 2^{M+1} + M - 1$, то $2^{2^M+1} \geq n - M + 1$ и $f_n(M+1) = 2^{2^M+1} + 2^{n-M-1} > 2^{2^M} + 2^{n-M+1}$.

12*) О. Б. Лупановым установлено, что $\lambda(n) \leq \frac{2^n}{n} (1 + \epsilon)$, и предложен соответствующий метод синтеза (Лупанов О. Б., О синтезе контактных схем, *ДАН*, 119, № 1, 1958, 23). Вместе с теоремой 7 (см. ниже) это дает асимптотику

$$\lambda(n) \sim \frac{2^n}{n}.$$

13*) а) 1°. Пусть $M + 2^M \leq n \leq 2^{M+1}$. Тогда $\frac{n}{2} \leq 2^M \leq n - M$

$$\text{и } \lambda(n) < 2(2^{n-M} + 2^{2^M}) \leq 4 \cdot 2^{n-M} = \frac{2^{n+2}}{2^M} \leq \frac{2^{n+3}}{n}.$$

2°. Пусть $2^{M+1} + 1 \leq n \leq 2^{M+1} + M$ (эти неравенства имеют смысл лишь при $M \geq 1$; тогда $n \geq 5$). Тогда $\frac{n-M}{2} \leq 2^M \leq \frac{n-1}{2}$ и

$$\lambda(n) < 2(2^{n-M} + 2^{2^M}) \leq 2(2^{n-M} + 2^{n/2}) \leq \\ \leq 2^{n+1} \left(\frac{1}{2^M} + \frac{1}{2^{n/2}} \right) \leq 2^{n+1} \left(\frac{2}{n-M} + \frac{1}{2^{n/2}} \right).$$

При $n \geq 4$ имеем $n \leq 2^{n/2}$, $M \leq \log_2 \frac{n-1}{2} \leq \frac{n}{3}$, $n-M \geq \frac{2}{3}n$ и $\frac{2}{n-M} + \frac{1}{2^{n/2}} \leq \frac{4}{n}$. Поэтому $\lambda(n) < \frac{2^{n+3}}{n}$

б) Пусть

$$2^M + 2M \leq n \leq 2^{M+1} - 2. \quad (**)$$

Тогда

$\lambda(n) < 2(2^{n-M} + 2^{2^M}) \leq 2(2^{n-M} + 2^{n-2M}) = 2^{n+1} \left(\frac{1}{2^M} + \frac{1}{2^{M^2}} \right)$. Так как $2^M \gg \frac{n+2}{2}$, то

$$\frac{1}{2^M} + \frac{1}{2^{M^2}} = \frac{1}{2^M} \left(1 + \frac{1}{2^M} \right) \leq \frac{2}{n+2} \cdot \frac{n+4}{n+2} = \frac{2(n+4)}{n^2+4n+4} < \frac{2}{n}.$$

Поэтому $\lambda(n) < \frac{2^{n+2}}{n}$ числа n , удовлетворяющие неравенству (**) при каких-нибудь M , составляют почти все натуральные числа.

в) Пусть $n_i = 2^i + 2i$. Тогда

$$\lambda(n_i) < 2(2^{n_i-i} + 2^{n_i-2i}) = \frac{2^{n_i+1}}{2^i} \left(1 + \frac{1}{2^i} \right) \sim \frac{2^{n_i+1}}{n_i} \quad (i \rightarrow \infty).$$

14*) Г. Н. Поваров предложил усовершенствование метода Шеннона — так называемый метод каскадов, позволяющий для отдельных функций строить простые схемы, а в некоторых случаях — даже минимальные (Поваров Г. Н., Математическая теория синтеза контактных (1, k)-полюсников, ДАН, 100, № 5, 1955, 909; Поваров Г. Н., Метод синтеза вычислительных и управляющих контактных схем, Автоматика и телемеханика, 18, № 2, 1957, 145).

Г. Н. Поваров исследовал также асимптотическое поведение функции $L(k, n)$ — минимального числа контактов, достаточного для реализации любой системы из k функций от n аргументов.

15*) См. сноску* на стр. 46.

16*) Здесь имеются в виду соединения полюсов, расположенных подряд, т. е.

$$(a, 1, 2, \dots, i_1)(i_1+1, \dots, i_2)(i_2+1, i_2+2, \dots) \dots (\dots, K, b).$$

17*) Э. Н. Гилберт (Gilbert E. N., *N-terminal switching circuits*, *Bell Syst. Techn. J.*, 30, № 3 (1951), 668) показал, что $f(k) > k^{(1-\varepsilon)k}$ (точнее, $f(k) > \left(c \frac{k}{\log^2 n}\right)^k$, где c — некоторая константа). Поэтому намеченный автором план доказательства не может быть осуществлен.

18*) О. Б. Лупановым было показано, что в случае параллельно-последовательных схем

$$\lambda(n) \sim \frac{2^n}{\log_2 n}$$

(Лупанов О. Б., О сложности реализации функций алгебры логики формулами, сб. Проблемы кибернетики, вып. 3 (1960), 61, а также Лупанов О. Б., О синтезе некоторых классов управляющих систем, сб. Проблемы кибернетики, вып. 10 (1963); во второй работе приводится доказательство непосредственно этого соотношения).

19*) См. прим. 12*).

20*) В частности, допускается перенос единицы от $a+1$ к a ($a \geq 2$), что равносильно перестановке этих чисел.

21*) Для последовательности $B = (b_1, b_2, \dots, b_n)$ пусть $S(B) = \sum_{i=1}^n i b_i$.

Более подробно доказательство достаточности можно провести по индукции с помощью следующего вспомогательного утверждения

Пусть монотонно неубывающие последовательности $B' = (b'_1, \dots, b'_n)$ и $B'' = (b''_1, \dots, b''_n)$ удовлетворяют условиям

$$(1') \sum_{i=1}^s b'_i \geq \sum_{i=1}^s b''_i, \quad s=1, 2, \dots, n,$$

$$(2') \sum_{i=1}^n b'_i = \sum_{i=1}^n b''_i$$

(3') B' и B'' имеют одинаковое число единиц.

Тогда если B' и B'' не совпадают, то в B'' можно осуществить допустимый перенос единицы, так что полученная последовательность B''' будет монотонно неубывающей и будут выполнены условия

$$(4') \sum_{i=1}^s b'_i \geq \sum_{i=1}^s b'''_i \geq \sum_{i=1}^s b''_i, \quad s=1, 2, \dots, n,$$

$$(5') S(B''') < S(B'').$$

В самом деле, пусть k — наименьшее число, такое, что $b'_k > b''_k$, а l — наибольшее число, такое, что $b'_l < b''_l$. Пусть m — наименьшее число, такое, что $b'_m < b''_m$ (из (2') следует, что такое m существует); ясно, что $m > l$ и $b'_m > b''_{m-1}$. Так как $b'_m \geq b'_l$, то $b'_m \geq b'_l + 2$. Последовательность B''' , полученная из B'' в результате переноса единицы от b''_m к b''_l , удовлетворяет всем перечисленным выше условиям.

22*) За исключением случая $n=3$. Но для этого случая все утверждение леммы проверяется непосредственно: разложение всех последовательностей $a_1, a_2, a_3 = (2, 4, 8)$ приведено на стр. 84.

23*) $i2^{q-1} \geq i2^{i+1}$. Если $i=1$, то $i2^{i+1}=4=4(2^i-1)$. Если $i \geq 2$, то $i2^{i+1} \geq 4 \cdot 2^i > 4(2^i-1)$.

24*) При $i \geq 3$. Рассмотрим теперь случай $i \leq 2$, $i=1$. В этом случае $q=2$, $3 \leq a_1 \leq 4$, $p \geq 3$, $a_2 \geq 2^2+1=5$. Если $a_1=3$, то $R=4$ и $\mu_1=6$; так как $a_1+a_2+a_3 \geq 2+4+8$ и $a_3 \geq a_2$, то $a_3 \geq 6$. Если $a_1=4$, то $R=3$, $\mu_1=5$; $a_3 \geq a_2 \geq 5$, $i=2$. В этом случае $p \geq 4$, $a_2 \geq 2^3+1=9$, $a_3+a_4 \geq 2a_2 \geq 18$, $\mu_1+\mu_2=3 \cdot 2^{i+1}+R-4=3 \cdot 2^{q-1}-2^{q-1}-4 \leq 18$, так как $q \geq 2$.

25*) Эта теорема доказывается так же, как теорема 7 (с использованием деревьев с почти равномерной нагрузкой).

26*) Точнее, $g(h(X_{i_1}, \dots, X_{i_s}), X_{i_{s+1}}, \dots, X_{i_n})$.

27*) Впоследствии был выделен еще ряд классов функций, допускающих более простую схемную реализацию, чем большинство функций. С. В. Яблонским построено и исследовано континуальное семейство классов функций, замкнутых относительно подстановок констант и переименования аргументов (Яблонский С. В., Об алгоритмических трудностях синтеза минимальных контактных схем, сб. Проблемы кибернетики, вып. 2, 1959, 75). Изучая эти классы, С. В. Яблонский установил, что построение «самых плохих» функций в некотором естественном классе алгоритмов связано с перебором всех функций алгебры логики.

ТРЕБОВАНИЯ, ПРЕДЪЯВЛЯЕМЫЕ К ОБЪЕМУ ПАМЯТИ ТЕЛЕФОННОГО КОММУТАТОРА¹⁾

1. Введение

Общий случай телефонного коммутатора с N абонентами изображен схематически на рис. 1. Основное назначение коммутатора заключается в установлении связи между любыми двумя абонентами. При этом коммутатор должен некоторым образом «помнить», какие абоненты связаны между собой, до тех пор пока не закончатся соответствующие разговоры. Для этого требуется определенное количество запоминающих элементов, зависящее от числа абонентов, максимальной частоты вызовов и т. д. На основе этих соображений ниже выводится ряд соотношений, дающих минимальное возможное число реле, переключателей системы кроссбар и других элементов, необходимых для выполнения этой функции запоминания. Сравнение характеристик любого предложенного проекта коммутатора с требованиями минимизации, полученными из данных соотношений, позволяет измерить эффективность применения данного проекта в терминах числа запоминающих элементов.

В физической системе память представляется посредством внутренних устойчивых состояний этой системы. Питание реле может осуществляться блокирующей схемой таким образом, что контакт будет оставаться замкнутым или разомкнутым в зависимости от его исходного состояния. Тогда реле имеет два устойчивых состояния. Набор из N реле имеет 2^N возможных комбинаций положений контактов и может быть соединен таким образом, чтобы все эти комбинации были устойчивыми.

В качестве меры памяти системы могло бы быть использовано полное число ее устойчивых состояний, но более удобно работать с логарифмом от этого числа. Основная причина этого состоит в том, что емкость памяти при такой мере пропорциональна числу используемых элементов. При N реле емкость памяти M равна $\log 2^N = N \log 2$. Если в качестве основания логарифмов взято число 2, то $\log_2 2 = 1$ и $M = N$. Получаемые в результате еди-

¹⁾ Shannon C., Memory requirements in a telephone exchange, *BSTJ*, 29 (1950), 343.

ницы измерения могут быть названы двоичными единицами, или битами.

Устройство с объемом памяти в M битов может сохранять M различных ответов «да» или «нет», т. е. M нулей или единиц. В ряде случаев также оказывается удобным использовать в качестве основания логарифмов число 10. Получаемые в результате единицы называются тогда десятичными единицами. Одно реле обладает памятью, имеющую емкость в 0,301 десятичных единиц. Переключатель системы кроссбар типа 10×10 содержит 100 гнезд. Если бы каждое гнездо действовало независимо от остальных, то общая емкость памяти составила бы 100 двоичных разрядов или 30,1 десятичных разрядов.



Рис. 1. Схема телефонного коммутатора.

Однако обычно в каждой вертикали может быть использовано только одно гнездо. При этом ограничении емкость памяти составляет один десятичный разряд на каждую вертикаль, т. е. всего десять десятичных разрядов. Панели, применяемые в телефонных коммутаторах панельного типа, представляют собой другую форму запоминающего устройства. Если число возможных уровней панельного переключателя равно 500, то емкость его памяти равна $\log 500$, т. е. 8,97 двоичных или 2,7 десятичных разрядов. Наконец, при шаговой системе применяются избирательные переключатели со ста гнездами. Эти переключатели обладают памятью в два десятичных разряда.

Часто оказывается, что фактически доступная емкость памяти группы реле или других устройств меньше, чем сумма емкостей отдельных устройств, вследствие искусственных ограничений, наложенных на допустимые состояния. В силу технических причин определенные состояния делаются недоступными: например, если реле A возбуждено, то реле B должно быть невозбужденным и т. д. В переключателе системы кроссбар нежелательно иметь более девяти работающих гнезд на одной горизонтали из-за скачкообразного увеличения нагрузки на кронштейн. Ограничения такого рода уменьшают емкость памяти, приходящуюся на один элемент, и влекут за собой невозможность выполнения минимальных требований, выводимых ниже.

2. Память, требуемая для S произвольных вызовов от N абонентов

Простейший случай имеет место, если предполагается, что коммутатор изолирован (не имеет связи с другими коммутаторами), и требуется, чтобы он мог обеспечивать любую возможную комбинацию из S или меньшего числа разговоров между абонентами. Если общее число абонентов равно N , то число способов выбора m пар дается формулой

$$\frac{N(N-1)(N-2)\dots(N-2m+1)}{2^m m!} = \frac{N!}{2^m m! (N-2m)!} \cdot \quad (1)$$

Числитель $N(N-1)\dots(N-2m+1)$ представляет собой число способов выбора $2m$ абонентов из N абонентов. Множитель $m!$ соответствует числу перестановок в порядке вызовов, а 2^m — числу инверсий абонентов в парах. Общее число возможных комбинаций равно сумме таких выражений для значений $m = 0, 1, \dots, S$, т. е.

$$\sum_{m=0}^S \frac{N!}{2^m m! (N-2m)!} \cdot \quad (2)$$

Для каждой из этих возможных комбинаций коммутатор должен обладать соответствующим устойчивым внутренним состоянием. Поэтому он должен обладать памятью емкости M , где

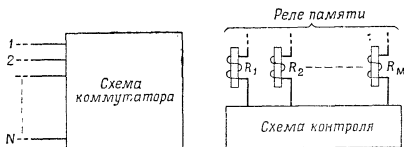
$$M = \log \sum_{m=0}^S \frac{N!}{2^m m! (N-2m)!} \cdot \quad (3)$$

Коммутатор, построенный только на реле, содержал бы по меньшей мере $\log_2 \sum N! / 2^m m! (N-2m)!$ реле.

Если переключатели системы кроссбар 10×10 применяются обычным способом, то коммутатор должен содержать по меньшей мере $\frac{1}{10} \log_{10} \sum N! / 2^m m! (N-2m)!$ этих переключателей и т. д. При использовании меньшего числа переключателей число устойчивых конфигураций связей не будет достаточным для того, чтобы различать все возможные желательные взаимосвязи абонентов. При $N = 10\,000$ и при максимальной нагрузке, например, в 1000 одновременных разговоров имеем M , равным 16 637 двоичным разрядам; при этом было бы необходимо по меньшей мере столько же реле или 502 переключателя системы кроссбар 10×10 . Между прочим, для чисел N и S такого порядка в выражении (3) существен только член, соответствующий значению $m = S$.

Подсчитанная выше емкость памяти соответствует только памяти, требуемой для выполнения основной функции запоминания разговаривающих абонентов на весь период разговора. При этом

не учитывались функции управления и контроля. Одну частную функцию управления легко учесть. Будем считать, что разговор ведет вызывающий абонент и связь разрывается в тот момент, когда он, а не вызванный им абонент вешает трубку. Таким образом, коммутатор должен различать случай, когда a звонит b , и случай, когда b звонит a . Вместо того чтобы подсчитывать число возможных пар, нам следует подсчитать число упорядоченных пар. В результате из приведенных выше формул удаляется множитель 2^m .



Р и с. 2. Коммутатор с минимальной памятью.

Возникает вопрос, являются ли эти пределы наилучшими возможными, т. е. возможно ли, например, построить коммутатор с применением только этого минимального числа реле. В принципе это возможно, но по различным причинам практически совершенно невыполнимо, если используются обычные типы реле или переключаательных элементов. На рис. 2 схематически изображен такой идеализированный коммутатор. Имеется M запоминающих реле, занумерованных числами $1, 2, \dots, M$. Каждая возможная конфигурация вызовов задается двоичным числом, заключенным между 0 и 2^M , и связана с соответствующей конфигурацией состояний реле. Имеется как раз столько таких состояний, сколько требуется для того, чтобы обеспечить все желаемые соединения абонентов.

Переключательная схема — это схема, состоящая из контактов запоминающих реле, причем если реле находятся в каком-либо данном состоянии, то нужные провода связаны между собой согласно заранее выбранному соответствию. Управляющая схема по существу является просто функциональным устройством и поэтому не требует памяти. После завершения разговора или в начале нового разговора желаемая конфигурация блокирующих реле сравнивается с имеющейся конфигурацией и напряжения подаются (или снимаются) на все реле, состояния которых должны быть изменены.

Нет нужды говорить, что коммутатор этого типа, хотя и использует минимальную память, имеет много недостатков, как это часто случается, когда минимизируется один параметр конструкции без учета других важных характеристик. В частности, относительно

схемы на рис. 2 можно заметить следующее: 1) каждое управляющее реле должно нести огромное количество контактов; 2) при каждом новом вызове или при завершении разговора должны измениться состояния многих запоминающих реле, что ведет к их быстрому износу и к помехам для идущих разговоров; 3) неисправность одного из запоминающих реле влечет за собой полный выход коммутатора из строя.

3. Условие раздельной памяти

Практическая невозможность построения коммутатора с абсолютно минимальной емкостью памяти наводит на мысль исследовать требования памяти при условии более реальных допущений. В частности, предположим, что при работе коммутатора за каждым происходящим в данное время разговором может быть закреплена отдельная часть памяти. В таком случае завершение текущего разговора или начало нового разговора не нарушит состояние элементов памяти, связанных с любым происходящим разговором. Это предположение достаточно хорошо выполняется при использовании стандартных типов коммутаторов и является естественным способом избежать трудностей 2) и 3), встречающихся при конструировании коммутаторов с применением абсолютно минимальной памяти.

Если коммутатор должен обеспечивать S одновременных разговоров, то должно быть по меньшей мере S отдельных запоминающих устройств. Кроме того, если их только S , то каждое ¹⁾ из этих запоминающих устройств должно обладать емкостью $\log 1/2N(N-1)$. Для того чтобы убедиться в этом, предположим, что все разговоры закончились, за исключением разговора, фиксируемого в некотором отдельном запоминающем устройстве. Тогда состояние всего коммутатора определяется состоянием этого отдельного запоминающего устройства. Зарегистрированный в этом устройстве разговор может происходить между любыми двумя из N абонентов, что дает всего $N(N-1)/2$ возможностей. Эти возможности должны соответствовать различным состояниям рассматриваемого запоминающего устройства, а следовательно, это устройство обладает емкостью, по меньшей мере равной

$$\log \frac{N(N-1)}{2}.$$

¹⁾ Б. Холбрук указал, что при применении более чем S запоминающих устройств каждое из этих устройств может при определенных значениях отношения $\frac{S}{N}$ обладать меньшей памятью, в результате чего получается чистая экономия. Однако это справедливо только для нереально больших частот вызовов.

В таком случае общая емкость требуемой памяти равна

$$M = S \log \frac{N(N-1)}{2}. \quad (4)$$

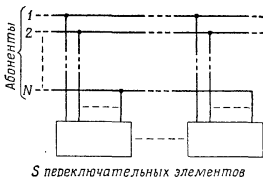
Если коммутатор должен помнить, какой из двух абонентов начал разговор, то получаем

$$M = S \log N(N-1), \quad (5)$$

или с высокой точностью при больших значениях N

$$M = 2S \log N. \quad (6)$$

Приближение, состоящее в замене выражения (5) на (6) и имеющее порядок $\frac{S}{N} \log e$, эквивалентно добавлению памяти, требуемой

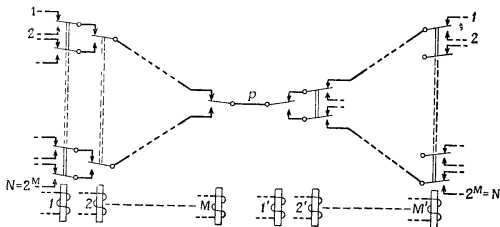


Р и с. 3. Коммутатор с минимальной раздельной памятью.

для того, чтобы допустить установление связей, идущих от абонента к нему самому. При $N = 10\,000$ и $S = 1\,000$ получаем из выражения (6) значение $M = 26\,600$. Значительное различие между минимальной требуемой памятью и памятью, фактически используемой в стандартных коммутаторах, частично обусловлено многими контролирующими и управляющими функциями коммутатора, которые не были учтены выше, а частично статистическими требованиями обеспечения резерва, компенсирующего ограниченные возможности скорости соединения абонентов.

Нижняя оценка, данная в (6), по существу реализуется в схематическом коммутаторе (рис. 3). Каждый блок содержит память емкостью $2 \log N$ и контактную схему, обеспечивающую связь любых двух входов, причем с каждым возможным состоянием памяти связана упорядоченная пара входов. На рис. 4 изображена такая схема. Посредством соответствующего возбуждения запоминающих реле $1, 2, \dots, M$ можно связать точку p с любым из $N = 2^m$ абонентов слева. Реле $1', 2', \dots, M'$ соединяют точку p с вызванным абонентом справа. Общая схема рис. 3 не слишком далека

от стандартной, хотя загрузка запоминающих элементов контактами является все же практически невозможной. В реальных панельных системах, системах кроссбар и шаговых системах доступ



Р и с. 4. Соединительная схема для коммутатора рис. 3.

блоков памяти к линиям ограничен для того, чтобы уменьшить загрузку контактами. Этим уменьшается гибкость системы связей, но статистически это уменьшение оказывается незначительным.

4. Связь с теорией информации

Формулу $M = 2S \log N$ можно интерпретировать в терминах теории информации¹⁾. Когда абонент поднимает трубку перед тем, как сделать вызов, результатом является выбор одной линии из множества, содержащего N линий. Если считать, что все абоненты могут начать разговор с одинаковой вероятностью, то соответствующее количество информации равно $\log N$. Когда абонент набирает нужный номер, происходит второй выбор одной из N возможностей. Общее количество информации, связанной с указанием исходного и конечного пункта вызова, равно $2 \log N$. Если возможно S одновременных разговоров, коммутатор должен помнить $2S \log N$ единиц информации.

Причина того, что этим методом получается формула «раздельной памяти», а не формула абсолютно минимальной памяти, состоит в том, что переоценена информация, содержащаяся в точном указании разговора. На самом деле вызывающий абонент должен принадлежать к числу незанятых; поэтому в общем случае происходит

¹⁾ Shannon C., A mathematical theory of communication. *Bell System Technical Journal*, 27, July, October 1948, 379, 623. [См. наст. сборник, стр. 243.—Прим. ред.]

выбор менее чем из N возможностей. Аналогично не может быть занят и вызванный участник разговора. Если вызванная линия занята, то разговор не может состояться и не требует памяти рассматриваемого здесь типа. При учете этих факторов получается формула абсолютно минимальной памяти. Условие раздельности памяти по существу эквивалентно предположению, что при запоминании очередного вызова коммутатор не использует информацию, которую он уже имеет в виде списка текущих разговоров.

В предположении, что все абоненты с равной вероятностью начинают разговор и с равной вероятностью вызывают любой номер, подсчет имеющейся информации соответствует в теории связи максимальной возможной информации или «энтропии». Если предположить вместо этого, что, как это и имеет место в действительности, определенные связи обладают высокой априорной вероятностью, тогда как для остальных вероятность относительно мала, то возможно произвести определенную статистическую экономию памяти.

В ограниченных пределах эта возможность уже используется. Предположим, что имеются два соседних района. Если вызов произведен из одного района, то вероятность того, что вызываемый абонент будет находиться в том же районе, значительно больше, чем вероятность его пребывания в другом районе. Таким образом, каждый коммутатор можно проектировать с таким расчетом, чтобы обслуживать местную телефонную связь и небольшое число межрайонных разговоров. В результате этого экономится память. Если у каждого коммутатора имеется по N абонентов и рассматривается предельный случай отсутствия связи между коммутаторами, то, согласно оценке (6), при этом общая память содержала бы $4S \log N$ двоичных разрядов, тогда как для обслуживания всех $2N$ абонентов одним коммутатором потребовалось бы $4S \log 2N$ двоичных разрядов.

Рассмотренная только что экономия возможна благодаря эффекту разбиения на группы. Существуют также статистики, затрагивающие особенности вызовов, осуществляемых отдельными абонентами. Девяносто процентов вызовов обычного абонента могут относиться к определенному небольшому числу людей, тогда как оставшиеся 10% возможных вызовов распределяются случайным образом между остальными абонентами. Это явление также может использоваться для уменьшения требуемой памяти, хотя составленные на бумаге проекты, включающие учет этого свойства, оказываются слишком сложными для практической реализации.

Краткое содержание

В работе исследуются реле, надежность которых можно описать в простых терминах с помощью теории вероятностей. Показывается, что, применяя надлежащим образом достаточно большое число таких реле, можно получить схемы, которые будут произвольно надежными независимо от степени надежности исходных реле. Описываются различные свойства таких схем.

Введение

В своей работе Дж. фон Нейман²⁾ рассматривает задачу построения надежных вычислительных схем из ненадежных элементов. Он исследует несколько случаев, один из которых, например, включает построение машин, использующих в качестве основного элемента элемент, реализующий штрих Шеффера³⁾.

Фон Нейман доказывает, что при определенных условиях можно скомбинировать ненадежные элементы штриха Шеффера так, чтобы получился элемент, который будет действовать подобно элементу штриха Шеффера высокой надежности. При известных условиях можно добиться надежного функционирования схемы за счет введения достаточно большого числа дополнительных элементов.

¹⁾ Moore E., Shannon C., Reliable circuits using less reliable relays, *Journal of the Franklin Institute*, № 3 (1956), 191; № 4, 281.

²⁾ von Neumann J., Probabilistic logic, California, Inst. of Technology, 1952. (Также опубликовано в «Automata Studies», ed. Shannon C., McCarthy J., Princeton University Press, 1956.) [Есть русский перевод, см. сб. Автоматы, ИЛ, 1956, Нейман Дж., Вероятностная логика и синтез надежных организмов из ненадежных компонент. — Прим. ред.]

³⁾ «Штрих Шеффера» это логическая функция двух переменных «не А и не В» ($\bar{A} \& \bar{B}$). Она обладает тем свойством, что через нее можно выразить все логические функции. «Элемент штриха Шеффера» — это элемент с двумя двоичными входами и одним двоичным выходом, который реализует эту логическую операцию. Ненадежный элемент такого типа дает правильную выходную величину только с некоторой вероятностью.

Настоящая статья возникла под влиянием работы фон Неймана и содержит аналогичный анализ для релейных схем. По-видимому, реле в основном более приемлемы для схем, исправляющих ошибки, чем нейроноподобные элементы, исследованные фон Нейманом. Во всяком случае наши результаты идут дальше в некоторых отношениях, чем результаты фон Неймана.

Прежде всего метод фон Неймана применим только при определенной, достаточно высокой надежности его элементов. При применении элемента штриха Шеффера абсолютно необходима вероятность ошибки меньшая чем $1/6$, а иногда при построении специальных схем, исправляющих ошибки, требуется вероятность ошибки порядка $1/100$ или еще меньшая. Развиваемые же нами методы применимы к произвольно ненадежным реле.

Далее, количество добавочных элементов, требуемое в наших схемах для заданного повышения надежности, значительно отличается от количества добавочных элементов, требующегося согласно фон Нейману. Например, в одном рассмотренном им численном примере требуется увеличение числа элементов приблизительно в 60 000 раз для получения определенного повышения надежности работы. Такое же повышение надежности достигается в релейных схемах при увеличении числа элементов только в 100 раз. В работе показано также, что в известном смысле некоторые из наших схем недалеко от минимальных. Полученные нами результаты говорят, в частности, о том, что в упомянутом численном примере необходимо иметь увеличение числа элементов по крайней мере в 67 раз для всякой схемы рассматриваемого типа. Следовательно, реальные схемы, для которых достигнуто требуемое повышение надежности за счет дополнительных элементов в количестве 100 к 1, не слишком плохи в смысле использования элементов.

Другое отличие заключается в том, что в случае использования реле нет необходимости применять то, что фон Нейман называет «восстанавливающей системой». При применении элементов фон Неймана окончательный результат (без восстановления) всегда имеет определенную остаточную ненадежность. При применении систем, описываемых здесь, эта вероятность ошибки может быть сделана сколь угодно малой.

Данная статья предназначена не для целей практического проектирования, а скорее для теоретического и математического изучения задачи. Однако возможны и некоторые практические применения. Надежность коммерческих реле обычно очень высока, например одна ошибка на 10^7 операций. Однако бывают случаи, когда даже такая надежность недостаточна.

Во-первых, в универсальных вычислительных машинах в ходе решения одной задачи срабатывает большое число реле, и ошибка в работе одного из них может вызвать ошибку в конечном резуль-

тате. Вследствие этого в вычислительных машинах фирмы «Белл» широко применяются самоконтролирующиеся устройства и устройства, исправляющие ошибки. Во-вторых, чрезвычайно высокая надежность требуется, когда от правильной работы релейной схемы зависит безопасность людей (железнодорожные блокировки, аварийные схемы на автоматических подъемниках, в управляемых снарядах и т. д.). Возможно, что некоторые из описанных нами простых схем можно будет использовать в таких случаях. Результаты, полученные в этой статье, непосредственно применимы к идеализированным реле, вероятность неверного срабатывания которых постоянна во времени, но не к техническим (реально существующим) реле, которые изнашиваются по мере их старения.

Идеализированные реле

Излагаемые результаты справедливы только для идеализированных реле, неисправность которых можно описать некоторым специальным способом в терминах вероятностей. В этих реле допускаются только случайные неисправности, причем только в предположении, что вероятность неисправностей остается постоянной с течением времени.

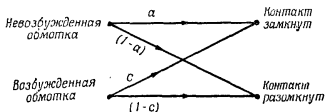
Наша идеализация не охватывает такие возможные в действительности случаи, когда реле изнашиваются от старения, когда обмотка реле перегорает или когда реле плохо припаяны. Предполагается также, что схема сконструирована и собрана правильно и что не может быть коротких замыканий между различными проходами.

Поскольку все вышеописанные типы ошибок и неисправностей могут иметь место на практике, результаты данной статьи непосредственно не применимы к техническим реле. Однако два типа неисправностей, рассматриваемых нами, в действительности имеют место в реле, так что предлагаемые схемы могут иметь некоторое применение.

Первый тип допустимых неисправностей — это неисправности в контакте, вследствие которых не осуществляется замыкание, что в технических реле часто происходит из-за наличия пыли. Второй тип неисправностей — это неразмыкание контакта, что в технических реле обычно происходит вследствие сваривающего действия тока, проходящего через контакты.

Будем рассматривать релейные схемы, в которых единственной причиной ошибок являются эти два типа неисправностей: контакт, который должен быть замкнут, остается разомкнутым, и контакт, который должен быть разомкнут, остается замкнутым. Предположим, что существуют две вероятности, относящиеся к контакту реле. Если реле не возбуждено, то контакт замкнут с вероятностью α и разомк-

нут с вероятностью $1 - a$. Если реле возбуждено, то контакт замкнут с вероятностью c и разомкнут с вероятностью $1 - c$; если a меньше c , то мы называем такой контакт *замыкающим*, если a больше c , то мы называем такой контакт *размыкающим*. Предположим, что различные контакты статистически независимы. В технических реле это, вероятно, не слишком далеко от истины для контактов на различных реле, а фактически больше ничего и не нужно для большинства результатов, которые будут получены. Кроме того, предположим, что в последующие моменты, когда обмотка реле снова возбуждается, «поведения» контактов статистически независимы.



Р и с. 1. Схематическое изображение переходных вероятностей.

Реле такого типа, управляемое вероятностями a и c , будем называть *ненадежным* реле. Вероятности его работы изображены схематически на рис. 1. Эта схема подобна тем, которые применяются для представления простого канала связи с шумами и действительно такое реле может рассматриваться как двоичный канал с шумом. Пропускная способность соответствующего канала будет тогда и только тогда равна нулю, когда $a = c$. Далее будет показано, что можно конструировать в высшей степени надежные вычислительные машины из большого числа ненадежных реле тогда и только тогда, когда $a \neq c$.

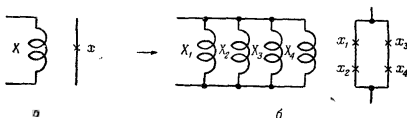
Общий метод повышения надежности

Общий метод повышения надежности, который будет описан, связан с конструированием схем, действующих подобно единичному контакту, но с большей надежностью, чем надежность контактов, из которых они состоят. Например, на рис. 2,а показано ненадежное реле X с замыкающим контактом x . Такое реле можно применять в качестве части большой вычислительной схемы. На рис. 2,б это реле заменено четырьмя ненадежными реле X_1, X_2, X_3, X_4 , обмотки которых, соединенные параллельно, заменяют одну обмотку X и контакты которых соединены параллельно-последовательно, причем эта двухполюсная схема заменяет один предыдущий контакт. Если вероятность замыкания каждого из этих четырех контактов

равна p , то ясно, что вероятность замыкания четырехконтактной схемы будет равна

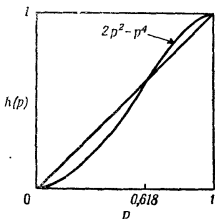
$$h(p) = 1 - (1 - p^2)^2 = 2p^2 - p^4.$$

График этой функции изображен на рис. 3. Эта кривая лежит выше диагонали $y = p$ для p , больших чем 0,618, и ниже ее для p , меньших чем 0,618. Это значит, что если 0,618 лежит между a и c (рис. 1),



Р и с. 2. Предлагаемый способ преобразования релейной схемы для повышения надежности.

то схема рис. 2,б будет действовать подобно реле с лучшими значениями, чем a и c , т. е. со значениями, более близкими соответственно к нулю и единице. Если, например, отдельные реле делали ошибки с вероятностями $a = 1 - c = 0,01$, то схема рис. 2,б будет



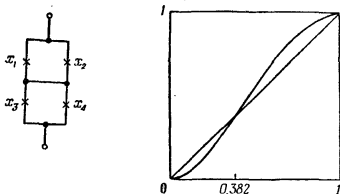
Р и с. 3. Функция, описывающая поведение схемы, изображенной на рис. 2,б.

делать ошибки, когда обмотка возбуждена с вероятностью 0,000396, а когда обмотка не возбуждена, — с вероятностью 0,0002. Таким образом, благодаря применению такой схемы достигается большое повышение надежности как в то время, когда обмотка возбуждена, так и в то время, когда обмотка не возбуждена.

На рис. 4 показано другое контактное устройство, которому соответствует иная функция

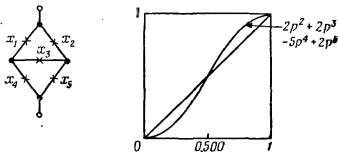
$$h(p) = [1 - (1 - p)^2]^2 = 4p^2 - 4p^3 + p^4.$$

Здесь опять $h(p)$ — вероятность замыкания схемы, когда каждый из отдельных контактов имеет вероятность замыкания p . Схема



Р и с. 4. Другая параллельно-последовательная схема и ее функция.

рис. 4 является двойственной к схеме рис. 2, а соответствующая кривая получается посредством замены нуля единицей и единицы нулем в абсциссе и в ординате кривой на рис. 3.



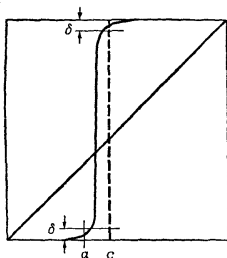
Р и с. 5. Мостиковая схема и ее функция.

Мостиковая схема рис. 5 дает симметричную кривую, пересекающую диагональ при $p = 0,5$. Для этой схемы имеем

$$h(p) = 2p^2 + 2p^3 - 5p^4 + 2p^5.$$

Все эти схемы стремятся приблизить p к его значениям 0 или 1 и таким образом повысить надежность. Далее будет показано, что многие другие схемы имеют подобные свойства. Кроме того, докажем, что для всякого положительного δ можно найти схему, кривая которой (рис. 6) пересекает диагональ при значении p , лежащем

между любыми двумя заданными числами a и c (независимо от того, насколько близки они друг к другу), причем значение функции меньше δ в точке a и больше $1 - \delta$ в точке c . Это значит, что из достаточного количества ненадежных реле можно получить сколь угодно надежное реле.



Р и с. 6. Общая форма кривой достижимых функций.

Ясно, что этот общий метод может применяться для повышения надежности как замыкающих, так и размыкающих контактов. Разница заключается только в обозначении точек a и c .

Свойства функции $h(p)$

Рассмотрим произвольную двухполюсную схему, построенную из контактов, вероятность замыкания каждого из которых равна p . Пусть при этом вероятность замыкания схемы равна $h(p)$. Исследуем некоторые из свойств функции $h(p)$.

Прежде всего, функция $h(p)$ есть многочлен; его можно записать следующим образом:

$$h(p) = \sum_{n=0}^m A_n p^n (1-p)^{m-n}, \quad (1)$$

где m — общее число контактов в схеме, A_n — число способов, которыми можно выбрать множество из n контактов в схеме так, что она будет замкнута, если эти n контактов будут замкнуты, а остальные контакты разомкнуты. Это очевидно, поскольку в правой части равенства (1) стоит сумма вероятностей различных способов, которыми схема может быть замкнута.

Первый ненулевой член в равенстве (1), скажем $A_s p^s (1-p)^{m-s}$ относится к кратчайшим путям через схему от одного полюса к другому; s есть длина этих путей, A_s — их число. Это объясняется тем, что в равенстве (1) все множества, соответствующие A_s , в действительности должны быть путями (иначе A_s не было бы первым ненулевым членом). Будем называть s *длиной* схемы. Из равенства (1) следует, что при значениях p , близких к нулю, функция $h(p)$ ведет себя как $A_1 p$.

Подобным образом можно поступить с вероятностью размыкания схемы и записать

$$1 - h(p) = \sum_{n=0}^m B_n (1-p)^n p^{m-n}, \quad (2)$$

где B_n — число множеств из n контактов, обладающих тем свойством, что если все контакты такого множества разомкнуты, а остальные контакты замкнуты, то схема будет разомкнута. Первый ненулевой член в этой сумме, скажем $B_t (1-p)^t p^{m-t}$, относится к наименьшему размыкающему множеству (т. е. множеству контактов, размыкание которых размыкает схему). Здесь t — число контактов в этом минимальном размыкающем множестве, B_t — число таких размыкающих множеств. Объяснение этого в основном такое же, как и выше. Назовем t *шириной* схемы. Очевидно, что при значениях p , близких к единице, $h(p)$ ведет себя как $1 - B_t (1-p)^t$.

Функция $h(p)$ может быть вычислена и другим способом. Пусть N — некоторый контакт в нашей схеме. Рассмотрим функцию вероятности $f(p)$ для схемы, полученной из исходной схемы посредством замыкания контакта N , и функцию вероятности $g(p)$ для схемы, полученной из исходной посредством размыкания этого контакта. Очевидно, что

$$h(p) = pf(p) + (1-p)g(p). \quad (3)$$

Кроме того, если $0 \leq p \leq 1$, то

$$f(p) \geq g(p). \quad (4)$$

Это интуитивно ясно из того, что любое замыкание в схеме не может уменьшить вероятность того, что схема будет замкнута. Формально это следует из равенства (1), если заметить, что случаи, когда схема, соответствующая функции g , замкнута, образуют подмножество множества случаев, в которых замкнута схема, соответствующая функции f , и, следовательно, члены в выражении (1) для функции f не меньше соответствующих членов в выражении для функции g .

Если данная схема плоская, то она будет иметь двойственную схему. Пусть $h_D(p)$ — функция вероятности для этой двойственной схемы. Пусть каждому состоянию контактов основной схемы соот-

ветствует состояние контактов двойственной схемы, при котором соответствующие контакты имеют противоположные состояния. Тогда состояния, для которых основная схема разомкнута, будут соответствовать состояниям, для которых двойственная схема замкнута. Если вероятность замыкания контактов в двойственной схеме равна $1 - p$, где p — вероятность замыкания контакта в основной схеме, то вероятности соответствующих состояний будут одинаковыми. Следовательно,

$$1 - h_D(1 - p) = h(p). \quad (6)$$

Пример такого соответствия между функциями h для некоторой схемы и её двойственной показан на рис. 3 и 4. Каждую из этих кривых можно получить из другой путем инвертирования, т. е. путем замены 0 на 1 и 1 на 0 на абсциссе и ординате.

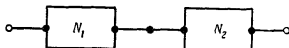
Если схема является самодвойственной (например, мостик рис. 5), то

$$1 - h(1 - p) = h(p). \quad (7)$$

Подставляя $p = 1/2$, находим $h(1/2) = 1/2$.

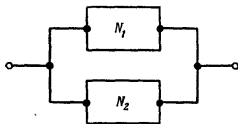
Комбинация двух схем

Рассмотрим теперь две схемы N_1 и N_2 с функциями $h_1(p)$ и $h_2(p)$. Если N_1 и N_2 соединяются последовательно (рис. 7).



$$h(p) = h_1(p)h_2(p)$$

Р и с. 7 Последовательное соединение двух схем.

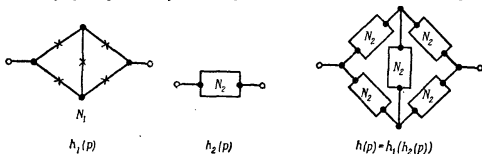


$$h(p) = 1 - (1 - h_1(p))(1 - h_2(p))$$

Р и с. 8. Параллельное соединение двух схем.

то полученная схема будет замкнутой тогда и только тогда, когда обе части будут замкнуты. Следовательно, функция $h(p)$ полученной схемы является произведением $h_1(p) \cdot h_2(p)$.

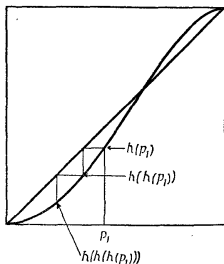
Если N_1 и N_2 соединяются параллельно (рис. 8), то полученная схема будет разомкнута тогда и только тогда, когда обе части будут разомкнуты; вероятность этого события равна



Р и с. 9. Композиция двух схем.

$(1 - h_1)(1 - h_2)$. Следовательно, функция $h(p)$ полученной параллельной схемы имеет вид $1 - (1 - h_1)(1 - h_2)$.

Третий способ комбинирования двух схем N_1 и N_2 осуществляется итерацией (подстановкой). Для этого нужно заменить каждый контакт в N_1 «копией» схемы N_2 . Типичный пример такой



Р и с. 10. Результат итерации композиций.

замены показан на рис. 9. Очевидно, что полученная схема имеет функцию h , являющуюся итерацией функций h двух исходных схем:

$$h(p) = h_1(h_2(p)). \quad (8)$$

Если N_1 и N_2 тождественны и этот процесс повторяется n 1—

раз, то получим n -ю итерацию функции h , т. е.

$$h^{(n)}(p) = h(h(h, \dots, h(p), \dots)).$$

Значение функции $h^{(n)}(p)$ можно определить по графику для $h(p)$ при помощи «ступенчатого» процесса, как показано на рис. 10 для $h^{(3)}(p_1)$. Таким образом, в схемах, графики которых пересекают диагональ только один раз, посредством подстановки можно достигнуть большого повышения надежности. Это повышение надежности в результате итераций, описываемое ступенчатым процессом рис. 10, очень сходно со случаями, описанными в работе фон Неймана.

Оценки функции $h'(p)$

Выведем одно интересное неравенство для производных возможных функций $h(p)$. Как следствие получим, что график всякой функции $h(p)$ может пересечь диагональ не более одного раза.

Теорема 1. Если $0 < p < 1$, то справедливо неравенство

$$\frac{h'(p)}{(1-h(p))h(p)} > \frac{1}{(1-p)p} \quad (9)$$

при условии, что $h(p)$ не равна тождественно нулю, единице или p .

Это утверждение будет доказано индукцией по числу контактов в схеме. Представим функцию $h(p)$ в виде (3), но при условии, что контакт, относительно которого осуществляется это представление, принадлежит некоторой цепи в схеме, и предположим, что для каждой из функций f и g либо доказываемое неравенство справедливо, либо эта функция является одной из трех исключительных функций; докажем неравенство для функции h . Поскольку рассматриваемый контакт принадлежит некоторой цепи, то из доказательства выражения (4) следует, что $f(p) > g(p)$ для всех p . Равенство $1 - f(p) + g(p) = 0$ не может выполняться ни для какого p , поскольку в противном случае $f(p) = 1$ и $g(p) = 0$, что означало бы, что нет пути через схему для функции g и нет размыкающего множества в схеме для функции f ; следовательно, $f(p) = 1$ и $g(p) = 0$ для всех p , т. е. $h(p) = p$, что противоречит условию теоремы.

Очевидно, что если $0 < p < 1$, то

$$(1-p)p(f-g)(1-f+g) > 0, \quad (10)$$

так как каждый из членов положителен. Раскрывая скобки, получаем

$$pf - pg - pf^2 + 2pfg - pg^2 - p^2f + p^2g + p^2f^2 - 2p^2fg + p^2g^2 > 0.$$

После некоторых преобразований имеем

$$-pf^2 + (1-p)pf - (1-p)g^2 - (1-p)pg > \\ > -[p^2f^2 + (1-p)^2g^2 + (1-p)2pfg].$$

Добавляя $pf + (1-p)g$ к обеим частям неравенства, находим

$$(1-f)pf + (1-p)pf + (1-p)(1-g)g - (1-p)pg > \\ > pf + (1-p)g - [pf + (1-p)g]^2 = h - h^2 = (1-h)h. \quad (11)$$

Поскольку по индуктивному предположению или $\frac{f'}{(1-f)f} > \frac{1}{(1-p)p}$, или f является одной из трех исключительных функций, в каждом случае $(1-f)f \leq (1-p)p f'$ и точно так же $(1-g)g \leq (1-p)p g'$. Из этих выражений и неравенства (11) получаем

$$(1-p)p^2f' + (1-p)pf + (1-p)^2pg' - (1-p)pg > (1-h)h.$$

Разделив на $(1-p)p$, находим

$$pf' + f + (1-p)g' - g > \frac{(1-h)h}{(1-p)p},$$

или

$$\frac{d}{dp}(pf + (1-p)g) > \frac{(1-h)h}{(1-p)p}, \\ \frac{h'}{(1-h)h} > \frac{1}{(1-p)p},$$

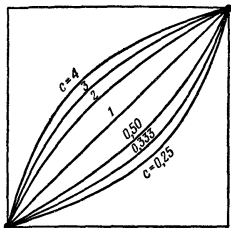
что и доказывает теорему.

Если заменить неравенство (9) в формулировке теоремы равенством, т. е. если положить $\frac{y'}{(1-y)y} = \frac{1}{(1-p)p}$, то получится дифференциальное уравнение, решения которого образуют однопараметрическое семейство функций. Из неравенства (9) следует, что допустимые функции h , соответствующие контактным схемам, должны иметь производные, большие чем производные функций этого семейства. Решая это дифференциальное уравнение, получаем

$$\frac{y(p)}{1-y(p)} = C \frac{p}{1-p}. \quad (12)$$

Графики функций этого семейства для $C = 1/4, 1/3, 1/2, 1, 2, 3, 4$ изображены на рис. 11. График всякой возможной функции $h(p)$ пересекает кривые этого семейства, так что в точке пересечения наклон графика функции $h(p)$ больше наклона кривой семейства. Следовательно, в открытом интервале $0 < p < 1$ график всякой функции $h(p)$ может пересечь каждую из кривых семейства самое большее один раз. Поскольку прямая линия с угловым коэффи-

циентом 1, проходящая через начало координат, есть одна из кривых этого семейства, то график произвольной функции $h(p)$ может пересекать эту линию самое большее один раз, скажем в точке $p = p_0$. Легко видеть [применяя ступенчатый процесс (рис. 10)], что функция $h^{(n)}(p)$ стремится к нулю для всех p , меньших p_0 , и стремится к единице для всех p , больших p_0 . Таким образом, если график



Р и с. 11. Семейство кривых, удовлетворяющих уравнению

$$\frac{y(p)}{1-y(p)} = C \frac{p}{1-p}.$$

функции $h(p)$ некоторой схемы пересекает диагональ, то итерации этой схемы имеют более высокую надежность. В пределе получаем

$$\lim_{n \rightarrow \infty} h^{(n)}(p) = \begin{cases} 1 & \text{при } p > p_0, \\ p_0 & \text{при } p = p_0, \\ 0 & \text{при } p < p_0, \end{cases}$$

где p_0 — абсцисса точки (обязательно единственной) пересечения кривой графика функций $h(p)$ с диагональю.

Можно получить оценку *сверху* для производной $h'(p)$, используя любопытным образом некоторые положения теории информации. Рассмотрим двоичный канал, изображенный на рис. 12. Скорость передачи информации для этого канала равна

$$\begin{aligned} R &= H(y) - H_z(y) = \\ &= -(p - \varepsilon Q) \log(p - \varepsilon Q) - (q + \varepsilon Q) \log(q + \varepsilon Q) + \\ &\quad + (1 - Q)(p \log p + q \log q) + \\ &\quad + Q[(p - \varepsilon) \log(p - \varepsilon) + (q + \varepsilon) \log(q + \varepsilon)]. \end{aligned}$$

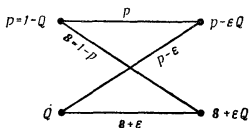
Для достаточно малых ε функция $(a + \varepsilon) \log(a + \varepsilon)$ может быть разложена в ряд Тейлора

$$a \log a + (1 + \log a) \varepsilon + \frac{\varepsilon^2}{a} + \dots$$

Используя это разложение в вышеприведенном выражении для всех членов, содержащих ε , получаем, что постоянный член и член первого порядка относительно ε обращаются в нуль. Первый ненулевой член определяется равенством

$$R = (Q - Q^2) \frac{\varepsilon^2}{pq} = \left[\frac{1}{4} - \left(Q - \frac{1}{2} \right)^2 \right] \frac{\varepsilon^2}{pq}.$$

Как видно из этого выражения, R достигает максимума (при изменении Q) при $Q = 1/2$. Это максимальное значение R является, по



Р и с. 12. Двоичный канал, используемый для получения верхней оценки для $h'(p)$.

определению, пропускной способностью C канала. Таким образом, при ε , близком к нулю, пропускная способность C (асимптотически) равна $\varepsilon^2/4pq$.

Рассмотрим теперь ненадежное реле, которое имеет вероятность замыкания p , когда обмотка возбуждена, и вероятность замыкания $p - \varepsilon$, когда обмотка не возбуждена. Это реле можно рассматривать как канал, для которого обмотка является входом, а контакт — выходом. Если ε достаточно мало, то пропускная способность реле равна $\varepsilon^2/4pq$. Если имеется n реле с одинаковыми p и ε , то общая пропускная способность системы, имеющей n обмоток в качестве входа и n контактов в качестве выхода, равна $n\varepsilon^2/4pq$, поскольку пропускная способность совокупности независимых каналов равна сумме их пропускных способностей.

Докажем, исходя из этих соображений о пропускной способности, что функция вероятности $h(p)$ для наших контактных схем удовлетворяет соотношению

$$\frac{dh}{dp} \leq \sqrt{\frac{n(1-h)h}{(1-p)p}}. \quad (13)$$

Рассмотрим схему N с n контактами и ее функцию вероятности $h(p)$. Пусть отдельные реле и контакты имеют параметры p_1 и ε (см. рис. 12). Тогда вся схема в целом будет действовать подобно одному реле с параметрами $h(p_1)$ и $h'(p_1)\varepsilon$ (если ε мало). Такое реле имеет пропускную способность $(h'\varepsilon)^2/4(1-h)h$. Она должна быть меньше или равна пропускной способности для случая, когда наши n реле используются наилучшим возможным способом. Следовательно

$$\frac{(h'\varepsilon)^2}{4(1-h)h} \leq \frac{n\varepsilon^2}{4(1-p_1)p_1}.$$

Это справедливо для всех p_1 , и после некоторых преобразований мы получим требуемый результат

$$h' \leq \sqrt{\frac{n(1-h)h}{(1-p_1)p_1}}.$$

Если это неравенство заменить равенством, то получим дифференциальное уравнение

$$\frac{\sqrt{n} dp}{\sqrt{(1-p)p}} = \frac{dh}{\sqrt{(1-h)h}}.$$

Решение этого уравнения имеет вид

$$\sqrt{n} \arcsin(1-2p) = \arcsin(1-2h) + \theta. \quad (14)$$

Для данного числа n контактов возможный график функции $h(p)$ должен пересекать кривые семейства в точках, где наклон графика $h(p)$ не больше наклона кривой семейства.

Другую верхнюю оценку функции $h(p)$, соответствующей схеме с n контактами, можно получить иным образом. Двухполюсная схема соответствует булевой функции n переменных. Однако невозможно реализовать все булевы функции, применяя только по одному замыкающему контакту для каждой переменной. Отвлечемся от этого условия реализуемости и рассмотрим класс всех булевых функций n переменных. Для каждой булевой функции определим функцию $h(p)$ или вероятность того, что функция равна единице при условии, что каждое переменное имеет (независимо) вероятность p быть равным единице. Возникает вопрос, какие булевы функции имеют функции вероятности $h(p)$ с максимальными производными и лучше всего повышают надежность схемы.

Булеву функцию n переменных назовем *функцией кворума*, если существует такое s ($0 \leq s \leq n$), что если в ней менее s переменных равно 1, то функция равна 0, а если более чем s переменных равно 1, то функция равна 1.

Теорема 2. Если график функции вероятности $h_Q(p)$ для произвольной функции кворума n переменных пересекает график

функции вероятности $h(p)$ некоторой другой булевой функции n переменных, то в точке пересечения p_0 справедливо неравенство

$$h'(p_0) < h'_Q(p_0),$$

т. е. функция кворума имеет максимальную производную. Кроме того,

$$h(p) > h_Q(p), \quad 0 < p < p_0,$$

$$h(p) < h_Q(p), \quad p_0 < p < 1.$$

Эта теорема утверждает, что в некоторых отношениях функции кворума являются лучшими из всех булевых функций для наших целей повышения надежности.

Доказательство. Для произвольной булевой функции n переменных функция $h(p)$, представленная в форме многочлена, является суммой членов вида $p^i q^{n-i}$; член такого вида соответствует каждому состоянию переменных, из которых i имеют значение 1, и при этом состоянии переменных функция также имеет значение 1. Функция кворума имеет значение 0 для всех состояний с i , меньшим s , и значение 1 для всех состояний с i , большим s . Следовательно, функция $h_Q(p)$ имеет вид

$$Ap^s q^{n-s} + \sum_{i=s+1}^n C_n^i p^i q^{n-i} \quad \text{при} \quad 0 \leq A \leq C_n^s.$$

Поскольку h не тождественно равна h_Q , но их значения совпадают при $p = p_0$, то в многочлене h будут отсутствовать некоторые члены с $i \geq s$ (по сравнению с h_Q) и многочлен h будет иметь некоторые дополнительные члены с $i \leq s$ (по сравнению с h_Q). Другими словами, можно написать

$$h(p) = \sum_{i=0}^n B_i p^i q^{n-i},$$

где $B_i \leq C_n^i$. Пусть $C(p) = ap^s q^{n-s} + \sum_{i=s+1}^n B_i p^i q^{n-i}$, где a — наименьшее из чисел B_s и A . Тогда

$$\begin{aligned} h_Q(p) &= C(p) + \sum_{i=\tau+1}^n D_i p^i q^{n-i}, \\ h(p) &= C(p) + \sum_{i=0}^{\tau} E_i p^i q^{n-i}, \end{aligned} \tag{15}$$

где D_i и E_i — неотрицательные целые числа, а τ равно $s-1$ или s в зависимости от того, что меньше, B_s или A .

Заметим, что для выражения вида $u(p) = p^i q^{n-i}$ имеет место равенство

$$u'(p) = ip^{i-1}q^{n-i} - (n-i)p^i q^{n-i-1} = \left(\frac{i}{p} - \frac{n-i}{q}\right) u(p) = \frac{i-pn}{pq} u(p).$$

Таким образом, $\frac{u'}{u} = \frac{i-pn}{pq}$ есть монотонно возрастающая функция i . Далее, все члены в сумме равенства (15) для h_Q соответствуют большим значениям i , чем значения в сумме для h . Если $u_Q(p)$ — произвольный член в сумме для h_Q и $u(p)$ — произвольный член в сумме для h , то

$$\frac{u'_Q}{u_Q} > \frac{u'}{u},$$

следовательно, существует такая постоянная K , что¹⁾

$$\frac{u'_Q}{u_Q} > K > \frac{u'}{u}$$

и

$$u'_Q > Ku_Q, \quad Ku > u'.$$

Суммируя первое неравенство по всем членам u_Q , а второе — по всем членам u , получим

$$\sum u'_Q > K \sum u_Q, \quad K \sum u > \sum u'.$$

Но при $p = p_0$ имеем $\sum u_Q = \sum u$, и, следовательно,

$$\sum u'_Q > \sum u',$$

$$h'_Q(p_0) > h'(p_0).$$

Вторую часть теоремы получаем непосредственно, так как если бы она была не верна, то, поскольку графики функций h и h_Q непрерывны, они должны были бы пересечься еще в некоторой точке, отличной от p_0 , что противоречило бы первой части теоремы.

Схемы заданной длины и ширины

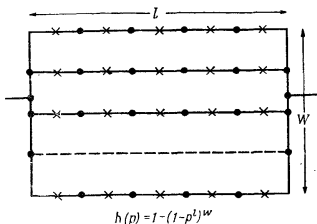
Было показано, что степени горизонтальности кривой $h(p)$ в окрестностях точек $p = 0$ и $p = 1$ связаны с длиной и шириной схемы. Ясно, что в практически важном случае интересующие нас значения p будут находиться в этих окрестностях, т. е. реле будут с самого

¹⁾ То есть число, не зависящее от i (но, вообще говоря, зависящее от p). — Прим. перев.

начала достаточно надежными. В этом параграфе будут приведены некоторые результаты о связи степени горизонтальности с числом элементов в схеме.

Теорема 3. Если схема N имеет длину l и ширину w , то она содержит по крайней мере lw контактов. Другими словами, если функция $h(p)$ ведет себя как Ap^l вблизи $p = 0$ и если $1 - h(p)$ ведет себя как $B(1 - p)^w$ вблизи $p = 1$, то соответствующая схема содержит по крайней мере lw контактов.

Доказательство. Сопоставим каждому контакту в N целое число следующим образом. Контактан, непосредственно



Р и с. 13. Параллельно-последовательная схема длины l и ширины w .

соединенным с левым полюсом схемы N , приписывается порядковый номер 1, контактам, непосредственно соединенным с контактами, обозначенными 1, приписывается 2 и т. д. Вообще контакту приписывается номер n , если существует цепь, идущая от него к левому полюсу и содержащая $n - 1$ других контактов, но не имеется такой цепи, содержащей меньшее число контактов.

Множество контактов, которым приписано число n , для любого фиксированного n от 1 до l образует размыкающее множество в схеме. В самом деле, каждая цепь через схему начинается у левого полюса контактом, обозначенным 1, и кончается у правого полюса контактом, обозначенным l или большим числом (если бы какие-либо из контактов, примыкающие к правому полюсу, были обозначены числами, меньшими l , то длина схемы N была бы меньше l). Вдоль каждой цепи при переходе от одного контакта к следующему номера изменяются на 0 или на ± 1 . Поэтому каждая цепь при переходе от контактов, занумерованных 1, к контактам, занумерованным числами, не превосходящими l , должна проходить через каждое промежуточное значение. Таким образом, если все контакты,

обозначенные n (для $1 \leq n \leq l$), будут исключены из N , то все цепи будут разорваны, и, следовательно, эти контакты образуют размыкающее множество.

Поскольку схема имеет ширину w , то каждое размыкающее множество будет содержать по крайней мере w контактов. Таким образом, по крайней мере w контактов будут обозначены 1, по крайней мере w контактов будут обозначены 2, ... и по крайней мере w контактов будут обозначены l . Следовательно, схема будет содержать по крайней мере lw контактов.

Другая формулировка теоремы 3 вытекает из замечаний, относящихся к уравнениям (1) и (2).

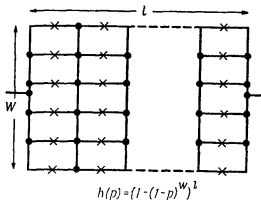


Рис. 14. Другой вариант параллельно-последовательной схемы длины l и ширины w .

Можно самыми разнообразными способами получить «размеры» l и w точно с lw контактами. Можно построить, например, последовательную схему с l контактами и соединить параллельно w таких схем (рис. 13). Можно поступить двойственным образом: w контактов соединить параллельно, а l таких схем — последовательно (рис. 14).

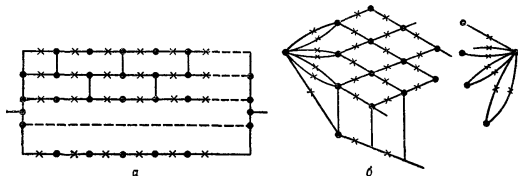
Теорема 4. Полная характеристика¹⁾ минимальных схем с размерами l и w состоит в следующем. Пусть Y и Z — полюсы, s_0 — множество, состоящее только из Y , и s_l — множество, состоящее только из Z . Кроме s_0 и s_l , имеются еще подмножества узлов $s_1, s_2, \dots, s_{l-1}$, такие, что; во-первых, ровно w контактов соединяют узлы из s_n с узлами из s_{n+1} ($n = 0, 1, \dots, l-1$) и, во-вторых, если узел из s_j имеет t элементов, соединяющих его с узлами из s_{j-1} , то он имеет t элементов, соединяющих его с узлами из s_{j+1} ($j = 1, 2, \dots, l-1$).

Это значит, что любую минимальную схему с размерами l и w можно получить из схемы рис. 13 посредством соответствующих сое-

¹⁾ То есть необходимое и достаточное условие. — Прим. перев.

динений (отождествлений) между узлами на одной и той же вертикальной линии. Если, например, все узлы на каждой вертикальной линии соединяются вместе, то в результате получится схема рис. 14. Другая возможная схема показана на рис. 15.

Чтобы показать, что всякая минимальная l - w -схема (т. е. схема с размерами l и w) имеет вид, описанный в формулировке теоремы 4, заметим сначала, что в предыдущем рассуждении каждое из занумерованных размыкающих множеств должно содержать ровно w элементов и эти элементы должны располагаться между элементами, имеющими меньшие номера, и элементами, имеющими

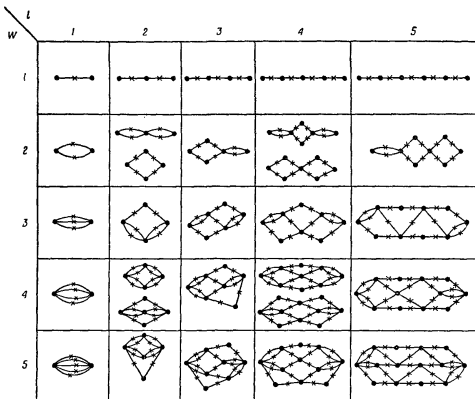


Р и с. 15. Гамакообразная схема длины l и ширины w .

бóльшие номера. Так, узлы между элементами с номерами $j-1$ и j будут относиться к подмножеству s_j . Предположим теперь, что какой-либо узел из s_i имеет m элементов, ведущих к узлам из s_{j-1} , и $m+p$ элементов, ведущих к узлам из s_{j+1} ($p>0$). Элементы с номерами $j+1$ образуют размыкающее множество из w элементов; легко видеть, что если эти $m+p$ элементов рассматриваемого размыкающего множества, исходящие из данного узла, заменить m элементами, идущими к узлам s_{j-1} , то получится размыкающее множество, имеющее меньше w элементов, что невозможно. Следовательно, всякая минимальная схема с размерами l и w есть схема рассматриваемого вида.

Чтобы доказать обратное, т. е. что всякая схема указанного вида имеет размеры l и w , заметим, во-первых, что при переходе от одного полюса к другому нужно пройти через узлы, принадлежащие $s_1, s_2, \dots, s_{l-1}$. Поэтому всякая цепь будет иметь длину по меньшей мере l и схема будет иметь длину l . Рассмотрим произвольное размыкающее множество c . Покажем, что c содержит по меньшей мере w контактов. Рассмотрим контакты в c , имеющие минимальные номера. Предположим, что один из них соединяет узел A из s_{j-1} с узлом B из s_i . Тогда или все элементы, идущие от B к узлам из s_{j-1} , принадлежат этому размыкающему множеству, или же какой-то из

них не принадлежит размыкающему множеству, и его можно устранить, получив при этом еще меньшее размыкающее множество. В первом случае эту группу элементов можно заменить таким же числом элементов, идущих от узла B к элементам из s_{j+1} с сохранением свойства быть размыкающим множеством. Поступая таким образом, будем постепенно передвигать размыкающее множество по направлению к правому полюсу, причем число элементов в множестве



Р и с. 16. Гамакообразные схемы различной длины и ширины.

будет или уменьшаться, или оставаться неизменным. Когда все элементы размыкающего множества окажутся рядом с правым полюсом, в этом множестве будет точно w контактов. Следовательно, в первоначальном размыкающем множестве их было не меньше, что и требовалось доказать.

Интересный тип минимальных l - w -схем получается при применении чередующихся соединений в схеме рис. 13, имеющих вид кирпичной кладки (рис. 15,а). «Стянув в точки» вертикальные соединения, получим схему, изображенную на рис. 15,б; схему такого типа будем называть *гамакообразной*. На рис. 16 показаны простые случаи гамакообразных схем. Ясно, что если оба числа l и w четные, то для них возможны две гамакообразные схемы. Если

же или l или w , или оба эти числа нечетные, то возможна только одна схема. Далее, двойственной схемой к гамакообразной схеме с длиной l и шириной w будет гамакообразная схема с длиной w и шириной l . Такие гамакообразные схемы являются в некотором смысле промежуточными между крайними минимальными l - w -схемами, изображенными на рис. 13 и 14, и имеют половину соединений, необходимых для перехода от схемы рис. 13 к схеме рис. 14. В случае когда l и w равны и нечетны, гамакообразная (единственная) схема будет самодвойственной.

Основная задача

Теперь приступим к решению нашей основной задачи — построению сколь угодно надежных схем из ненадежных реле. Покажем, что может быть построена схема с функцией $h(p)$, график которой изображен на рис. 6: $h(p)$ будет возрастать от значения, меньшего δ , при $p = a$, до значения, большего $1 - \delta$, при $p = c$, для любых $\delta > 0$, $0 \leq a < c \leq 1$. Оценим также число контактов в такой схеме в зависимости от значений δ , a и c .

Задача в общем случае разбивается на три этапа. Назовем их начальным, средним и последним.

Начальный этап состоит в нахождении схемы, которая, образно говоря, «передвигает» a и c так, чтобы они оказались по разные стороны от точки $p = 1/2$, т. е. для которой $h(a) < 1/2$ и $h(c) > 1/2$. Эту схему можно себе представить как единичный контакт с более хорошими значениями a и c (по крайней мере более приемлемыми для применяемого нами метода).

Средний этап состоит в построении схемы, которая раздвигает значения a и c из окрестности точки $1/2$, так чтобы a оказалось близким к $1/4$, а c — к $3/4$.

Последний этап заключается в построении схемы, осуществляющей передвижение этих точек из окрестностей $1/4$ и $3/4$ соответственно к 0 и 1.

Решение задачи в общем случае основано на замещении копией первой (полученной на начальном этапе) схемы каждого контакта второй (полученной на среднем этапе) схемы. Копией получившейся схемы заменяется каждый контакт схемы, полученной на последнем этапе. Тогда общее число контактов равно произведению чисел контактов в указанных трех вспомогательных схемах.

Разделение задачи на три части сделано главным образом для удобства математического описания. По-видимому, наиболее эффективный метод синтеза часто не будет состоять из такой итерации трех схем.

Во многих случаях, конечно, первая часть или даже две первых части решения не являются необходимыми, если a и c уже разде-

лены точкой $1/2$ и достаточно удалены друг от друга. Это всегда имеет место для реле, используемых на практике. Общий случай представляет главным образом теоретический интерес. Для практического случая, когда a и c уже близки к 0 и 1, будет показано, что схемы, предлагаемые на последнем этапе, достаточно экономны в смысле числа контактов. Они содержат не намного больше контактов, чем требует получаемая здесь нижняя оценка для всех схем.

Начальный этап

В первую очередь покажем, что можно построить схему, график функции $h(p)$ которой пересекает прямую $h(p) = 1/2$ в произвольном интервале и имеет средний наклон в этом интервале не менее $1/2$. Затем оценим число контактов в такой схеме. Точнее, докажем следующее.

Теорема 5. Пусть a и c таковы, что $0 < a < c < 1$, пусть далее $b = (a + c)/2$, $d = \max(b, 1 - b)$, $\varepsilon = (c - a)/4$. Тогда существует схема N , имеющая не более чем $\left[\frac{\log \varepsilon}{\log d} \right]$ контактов и такая, что $h_N(a) \leq \frac{1}{2} - \varepsilon$, $h_N(c) \geq \frac{1}{2} + \varepsilon$.

Лемма 1. Существуют две последовательности схем $N_0, N_1, N_2, \dots$ и $M_0, M_1, M_2, \dots$, такие, что для каждого i

- 1) N_i имеет не более i контактов;
- 2) M_i имеет не более i контактов;
- 3) $h_{N_i}(b) < \frac{1}{2} \leq h_{M_i}(b)$;
- 4) $h_{M_i}(b) - h_{N_i}(b) \leq d^i$;
- 5) либо M_i может быть получена соединением накоротко некоторых двух узлов схемы N_i , либо N_i может быть получена размыканием некоторого контакта схемы M_i .

Схемы M_i и N_i этой леммы получаются из «лестничной» схемы, общий вид которой изображен на рис. 17 (однако с различными числами горизонтальных и вертикальных элементов в соответствии

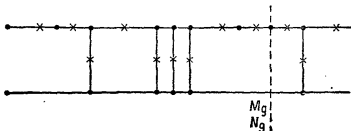


Рис. 17. Лестничная схема, использованная для доказательства теоремы 5.

со значением b и, возможно, начинающейся не горизонтальной, а вертикальной группой). Схемы M_i и N_i образуются при отрезании части схемы позади i -го контакта (как показано для M_0 и N_0) и соответственно замыкания или размыкания образовавшихся концов. Можно считать, что график функции $h(p)$ бесконечной схемы пересекает $h(p) = 1/2$ точно при $p = b$. Схемы M_i и N_i являются конечными аппроксимациями, которым соответствуют точки пересечения, лежащие левее и правее b .

Доказательство. Пусть M_0 — тождественно замкнутая схема и N_0 — тождественно разомкнутая схема. Эти схемы удовлетворяют условиям леммы для случая $i = 0$, так как $h_{M_0}(p) = 1$, $h_{N_0}(p) = 0$. Предположим теперь, что условия леммы выполнены для $i - 1$. Если M_{i-1} получается посредством соединения накоротко двух узлов схемы N_{i-1} , то пусть M — схема, получающаяся из N_{i-1} введением одного контакта между этими узлами. Если N_{i-1} получается путем размыкания одного из контактов в M_{i-1} , то пусть M — схема, получающаяся в результате введения нового контакта последовательно с этим контактом. Тогда M имеет не более i контактов; если замкнуть (разомкнуть) добавленный контакт в схеме M , то получится M_{i-1} (N_{i-1})¹⁾. Тогда на основании равенства (3)

$$h_M(p) = ph_{M_{i-1}}(p) + (1-p)h_{N_{i-1}}(p),$$

и, следовательно $h_{N_{i-1}}(p) < h_M(p) < h_{M_{i-1}}(p)$. Если $h_M(b) < \frac{1}{2}$, полагаем $N_i = M$ и $M_i = M_{i-1}$. Тогда

$$\begin{aligned} h_{M_i}(b) - h_{N_i}(b) &= h_{M_{i-1}}(b) - h_M(b) = \\ &= h_{M_{i-1}}(b) - [bh_{M_{i-1}}(b) + (1-b)h_{N_{i-1}}(b)] = \\ &= (1-b)[h_{M_{i-1}}(b) - h_{N_{i-1}}(b)] \leq \\ &\leq (1-b)d^{i-1} \leq d^i. \end{aligned}$$

Если же $h_M(b) \geq \frac{1}{2}$, то полагаем $M_i = M$ и $N_i = N_{i-1}$. Тогда, точно так же $h_{M_i}(b) - h_{N_i}(b) \leq bd^{i-1} \leq d^i$. Лемма доказана.

Лемма 2. $h'(p) \geq 3/4$ для всех p , таких, что

$$1/2 - \varepsilon \leq h(p) \leq 1/2 + \varepsilon.$$

В самом деле, так как $\varepsilon = \frac{c-a}{4} \leq \frac{1}{4}$, то $\frac{1}{4} \leq h(p) \leq \frac{3}{4}$; следовательно, $h'(p) \geq \frac{(h-1)h}{(p-1)p} \geq \frac{3/16}{1/4} = \frac{3}{4}$.

¹⁾ Или схема, реализующая ту же функцию. — Прим. перев.

Для доказательства теоремы 5 положим $i = \left[\frac{\log \varepsilon}{\log d} \right]$; тогда $d^i \leq \varepsilon$ и, следовательно,

$$h_{N_i}(b) < \frac{1}{2} \leq h_{M_i}(b) \leq h_{N_i}(b) + \varepsilon.$$

Пусть N — та из схем N_i и M_i , которая удовлетворяет условию $\left| h_N(b) - \frac{1}{2} \right| \leq \frac{\varepsilon}{2}$ (поскольку одна из них должна удовлетворять этому условию).

Тогда без ограничения общности достаточно показать, что

$$h_N(c) \geq \frac{1}{2} + \varepsilon.$$

Допустим противное, т. е. допустим, что $h_N(c) < \frac{1}{2} + \varepsilon$.

Тогда, поскольку функция h_N монотонна, то $\frac{1}{2} - \varepsilon \leq h(p) \leq \frac{1}{2} + \varepsilon$ для всех p , заключенных между b и c . Следовательно, по лемме 2 в этом интервале $h'(p) \geq \frac{3}{4}$. Поэтому

$$h(c) = h(b) + \int_b^c h'(p) dp \geq \frac{1}{2} - \frac{\varepsilon}{2} + \frac{3}{4}(c-b) \geq \frac{1}{2} - \frac{\varepsilon}{2} + \frac{3\varepsilon}{2} \geq \frac{1}{2} + \varepsilon,$$

что противоречит предположению.

Средний этап

Второй этап решения нашей задачи состоит в нахождении схемы, которая передвигает вероятности, немного меньшие и немного большие $1/2$, к значениям, меньшим $1/4$ и большим $3/4$ соответственно. Иными словами, найдем схемы, для которых

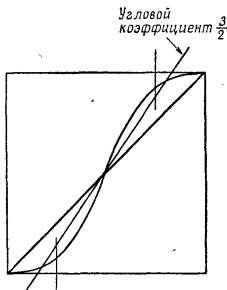
$$h\left(\frac{1}{2} - \varepsilon\right) \leq 1/4,$$

$$h\left(\frac{1}{2} + \varepsilon\right) \geq 3/4,$$

и подсчитаем число используемых в них контактов в зависимости от значения ε .

Схемы, удовлетворяющие нашим условиям, могут быть получены итерацией некоторой самодвойственной схемы с нею же самой достаточное число раз. Например, можно использовать гамакообразную схему три-на-три рис. 16. Нетрудно показать, что график функции $h(p)$ этой схемы, как показано на рис. 18, расположен ниже прямой с наклоном $3/2$, проходящей через точку $(1/2, 1/2)$, в интервале $1/4 < p \leq 1/2$ и выше этой прямой в интервале $1/2 < p \leq 3/4$. Отсюда видно, что итерация этой гамакообразной схемы с собой

достаточное число раз, обеспечивающая перемещение точки $\frac{1}{2} - \varepsilon$ к точке $\frac{1}{4}$, требует меньшего числа подстановок, чем число шагов в ступенчатом процессе при соответствующем движении вдоль этой прямой линии. Далее, если перемещаться вдоль прямой от точки $\frac{1}{2} - \varepsilon$, то очевидно, что после s итераций точка дойдет до значения $\frac{1}{2} - \varepsilon \left(\frac{3}{2}\right)^s$. Если желательно дойти до $\frac{1}{4}$ «по прямой»,



Р и с. 18. Функция, связанная с гамакообразной схемой три-на-три.

то s итераций будет достаточно в том случае, когда s есть минимальное число, удовлетворяющее условию

$$\frac{1}{4\varepsilon} \leq \left(\frac{3}{2}\right)^s.$$

Этого числа итераций будет заведомо достаточно для гамакообразной схемы, так как ее действие в ступенчатом процессе сильнее, чем в случае прямой линии. Число контактов в s -кратной итерации этой гамакообразной схемы равно

$$N_2 = 9^s \leq 9 \left(\frac{1}{4\varepsilon}\right)^{\frac{\log 9}{\log 3/2}} \leq 9 \left(\frac{1}{4\varepsilon}\right)^{5,41}.$$

Те же самые рассуждения применимы, конечно, для передвижения $\frac{1}{2} + \varepsilon$ к $p = \frac{3}{4}$, и это осуществляется той же самой схемой, поскольку она является самодвойственной.

Так как наше ε на начальном шаге определялось как $\frac{c-a}{4}$, то схема второго этапа требует самое большее $9 \left(\frac{1}{c-a} \right)^{5,41}$ контактов. Общая схема (после первого и второго этапов) требует не более чем $9 \left[\frac{\log \frac{c-a}{4}}{-\log d} \right] \left(\frac{1}{c-a} \right)^{5,41}$ контактов.

Последний этап

Наша третья схема последнего этапа должна передвигать вероятности $1/4$ и $3/4$, или лучшие, к точкам, близким к 0 и 1 соответственно. Этот тип схем также основан на итерации самодвойственной схемы, в частности гамакообразной схемы три-на-три или пятиэлементного мостика. Сначала будет удобно произвести подсчет контактов, с помощью которых возможно приблизиться к 0 и 1 в рассматриваемом случае.

В некоторых случаях возможно найти верхнюю и нижнюю оценки функции $h(p)$ в зависимости от p . Например, пусть

$$h(p) \leq ap^r. \quad (16)$$

Если составляется итерация соответствующей схемы, то мы получаем, используя тот факт, что $h(p)$ монотонна,

$$\begin{aligned} h^{(2)}(p) &\leq a (ap^r)^r = a^{1+r} p^{r^2}, \\ h^{(3)}(p) &\leq a^{1+r} (ap^r)^{r^2} = a^{1+r+r^2} p^{r^3}, \\ &\dots \dots \dots \\ h^{(s)}(p) &\leq a^{1+r+r^2+\dots+r^{s-1}} p^{r^s} = \\ &= a^{\frac{r^s-1}{r-1}} p^{r^s} = \frac{(a^{\frac{1}{r-1}})^{r^s}}{a^{\frac{1}{r-1}}} p^{r^s}. \end{aligned} \quad (17)$$

Теперь предположим, что исходная схема содержит n_1 контактов. Тогда ее s -кратная итерация с собой будет содержать $n = n_1^s$ контактов. Исключая s из выражения (17), получаем

$$h^{(s)}(p) \leq \frac{(a^{\frac{1}{r-1}})^n \frac{\log r}{\log n_1} p^{\frac{n \log r}{\log n_1}}}{\frac{1}{a^{\frac{1}{r-1}}}}.$$

Если бы знак в неравенстве (16) был обращен в другую сторону, то результат получился бы также обратным. Следовательно, будет справедлив аналогичный результат, если p заменить на $1-p$.

Гамакообразная схема три-на-три, рис. 16, имеет функцию вероятности

$$\begin{aligned} h(p) &= 8p^3 - 6p^4 - 6p^5 + 12p^7 - 9p^8 + 2p^9 = \\ &= 8p^3 - p^4 [6 + 6p - 12p^3 + 9p^4 - 2p^5] = \\ &= 8p^3 - p^4 [6(1 - p^3) + 6p(1 - p^2) + 7p^4 + 2p^4(1 - p)]. \end{aligned}$$

Так как выражение в квадратных скобках, очевидно, неотрицательно, то $h(p) \leq 8p^3$ для $0 \leq p \leq 1$. Следовательно, для s -кратной итерации в силу (17)

$$h^{(s)}(p) \leq \frac{1}{\sqrt[3]{8}} (\sqrt[3]{8} p)^{3s}.$$

Так как каждая итерация умножает число контактов на 9, то общее их число равно $N = 9^s$, $\sqrt[3]{N} = 3^s$ и

$$h^{(s)}(p) \leq \frac{1}{\sqrt[3]{8}} (\sqrt[3]{8} p)^{\sqrt[3]{N}}, \quad (18)$$

$$N \leq \left(\frac{\log \sqrt[3]{8} h^{(s)}(p)}{\log \sqrt[3]{8} p} \right)^2. \quad (19)$$

Следует заметить, что данный результат справедлив только при $p \leq \frac{1}{\sqrt[3]{8}} \approx 0,353$. Для больших значений p верхняя оценка h превосходит 1.

Обратим еще внимание на то, что в описанном процессе N растет скачкообразно, увеличиваясь при каждом скачке в 9 раз. Если правая часть неравенства (19) дана, то может потребоваться число N контактов, в 9 раз большее, чем определяемое неравенством, но этого количества будет заведомо достаточно.

Двойственная ситуация дает следующее неравенство для N :

$$N \leq \left[\frac{\log \sqrt[3]{8} (1-h)^{(s)} (1-p)}{\log \sqrt[3]{8} (1-p)} \right]^2. \quad (20)$$

Вернемся теперь к задаче последнего этапа. Пусть требуется улучшить вероятность от $p \leq \frac{1}{4}$ до, скажем, $h(p) \leq \delta$. Это требует, согласно неравенству (19), не более чем $9 \left(\frac{\log \sqrt[3]{8} \delta}{\log \sqrt[3]{8}} \right)^2$ контактов. Так как используется гамакообразная схема, являющаяся самодвойственной, то те же самые построения применимы, чтобы улучшить вероятность от $\frac{3}{4}$ до $1 - \delta$. Это делает та же самая схема.

Можно суммировать рассуждения, касающиеся итерации этих трех схем, следующим образом.

Теорема 6. Пусть даны $\delta > 0$, $0 < a < c < 1$ и пусть $d = \max\left(\frac{a+c}{2}, 1 - \frac{a+c}{2}\right)$. Тогда существует такая схема, что $h(a) < \delta$, $h(c) > 1 - \delta$, содержащая не более N контактов, где

$$N = 81 \left[\frac{\log \frac{c-a}{4}}{-\log d} \right] \left(\frac{1}{c-a} \right)^{\frac{\log}{\log 3/2}} \left(\frac{\log \sqrt{8}\delta}{\log \sqrt{8}} \right)^2.$$

Это выражение показывает, что эквивалент для сколь угодно надежного реле может быть получен из достаточного числа произвольно ненадежных реле, и дает оценку необходимого числа реле. Эта оценка безусловно не является вполне точной. В частности, множитель 81, который был введен («из соображений четности»), по-видимому, может быть уменьшен, если произвести более точный анализ. Для отдельных значений c , a и δ этот множитель будет равен единице, а в случаях, где множитель больше, подстановка других схем в гамакообразную схему три-на-три часто уменьшает этот множитель до числа, близкого к единице без значительного изменения других членов правой части равенства.

Оценки числа контактов в случае, когда реле уже достаточно надежны

Если реле с самого начала являются достаточно надежными, т. е. значения a и c соответственно меньше $1/4$ и больше $3/4$, то первые два шага описанного выше процесса могут быть опущены. Это, конечно, наиболее часто встречающийся на практике случай. Если он имеет место и требуется уменьшить значение a до δ_1 , а значение c увеличить до $1 - \delta_2$, то достаточно не более чем

$$\max \left[9 \left(\frac{\log \sqrt{8}\delta_1}{\log \sqrt{8}a} \right)^2, 9 \left(\frac{\log \sqrt{8}\delta_2}{\log \sqrt{8}(1-c)} \right)^2 \right]$$

контактов, и если числа δ_1 и δ_2 выбраны удачно, то множитель 9 может быть заменен, как правило, единицей.

В приложении приведены некоторые другие оценки, получаемые использованием гамакообразных схем общего вида. Эти оценки несколько более точны, поскольку в них, с одной стороны, устранен множитель 9 и, с другой стороны, коэффициент $\sqrt{8}$ заменен меньшим.

Теперь выведем неравенство, дающее нижнюю оценку числа контактов, необходимого для заданного повышения надежности.

Теорема 7. Пусть $0 < a < c < 1$ и N — двухполюсная схема с функцией вероятности $h(p)$, удовлетворяющей условиям

$$\begin{aligned} h(a) &\leq \delta_1, \\ h(c) &\geq 1 - \delta_2. \end{aligned}$$

Тогда число n контактов в этой схеме удовлетворяет условию

$$n \geq \frac{\log \delta_1}{\log a} \cdot \frac{\log \delta_2}{\log (1-c)}.$$

Например, если контакты не срабатывают один раз из десяти как в случае, когда они должны быть разомкнуты, так и в случае, когда они должны быть замкнуты, то будем иметь $a = 1 - c = 10^{-1}$. Если схема должна делать не более одной ошибки на 10^6 операций, то, согласно этой теореме, она должна иметь контактов не менее чем

$$\frac{\log 10^{-6}}{\log 10^{-1}} \cdot \frac{\log 10^{-6}}{\log 10^{-1}} = 36.$$

Для доказательства теоремы предположим, что схема N имеет длину l и ширину w . Тогда в ней имеется путь, проходящий через l контактов от одного полюса к другому. Вероятность того, что этот путь замкнут при условии, что все реле не возбуждены, равна a^l . Если это имеет место, то схема, конечно, будет замкнута (при условии, что все реле не возбуждены). Следовательно, $\delta_1 \geq a^l$ и $\log \delta_1 \geq l \log a$. Разделив обе части неравенства на отрицательное число $\log a$, получим

$$\frac{\log \delta_1}{\log a} \leq l.$$

Аналогичные рассуждения по отношению к размыкающему множеству из w элементов дают

$$\frac{\log \delta_2}{\log (1-c)} \leq w.$$

Так как все части обоих последних неравенств положительны, то можно эти неравенства перемножить. Получаем, используя неравенство теоремы 3,

$$\frac{\log \delta_1}{\log a} \cdot \frac{\log \delta_2}{\log (1-c)} \leq lw \leq n.$$

Резюмируя, можно сказать, что число n контактов, необходимое для уменьшения вероятности ошибки замыкания от a до δ_1 и вероятности ошибки размыкания от $1 - c$ до δ_2 , должно быть порядка

$$\frac{\log \delta_1}{\log a} \cdot \frac{\log \delta_2}{\log (1-c)}.$$

Это число не может быть меньше найденной оценки, и для некоторой бесконечно возрастающей последовательности значений n оно лишь немного больше, как это показано в приложении для некоторых гамакообразных схем.

Сравнение с элементами фон Неймана

В качестве численного примера можно рассмотреть случай, аналогичный разобранным в работе фон Неймана, в котором из элементов штриха Шеффера с вероятностью ошибки $1/200$ строятся элементы штриха Шеффера с вероятностью ошибки порядка 10^{-20} . Фон Нейман установил, что его схема требует примерно 60 000 элементов для получения требуемой надежности. Оказывается также, что число 60 000 слабо зависит от окончательной вероятности 10^{-20} и изменяется от 32 000 до 69 000, когда окончательная вероятность меняется от 10^{-17} до 10^{-23} .

В качестве близкого примера рассмотрим реле с исходными вероятностями ошибок $a = 1 - c = 1/200$ и найдем схему, которая понижает эти вероятности примерно до 10^{-20} . Так как наши исходные вероятности уже относительно хороши, необходимо проделать только последний этап.

Сначала используем итерации только гамакообразных схем три-на-три. Если используется одна ступень, то число контактов увеличивается в 9 раз и окончательная вероятность ошибки будет меньше чем $8(1/200)^3 = 10^{-6}$. Если используются две ступени, то число контактов увеличивается в 81 раз и окончательная вероятность ошибки будет меньше чем $8^{-1/2}(\sqrt{8} \cdot 1/200)^9 = 8 \cdot 10^{-18}$. Таким образом, при увеличении числа контактов в 81 раз получаем в общем надежность, близкую к желаемой, но, правда, немного меньшую. Следующая итерация схем три-на-три увеличивает число контактов в 729 раз и дает окончательную вероятность ошибки меньше чем $4 \cdot 10^{-51}(!)$.

Используя большие гамакообразные схемы (см. в приложении об оценках вероятностей ошибок), можно решить нашу задачу более точно. В частности, гамакообразная схема десять-на-десять дает вероятность ошибки, меньшую чем $2 \cdot 10^{-19}$, при увеличении числа контактов в 100 раз; схема одиннадцать-на-одиннадцать дает $2,2 \cdot 10^{-21}$ при увеличении в 121 раз.

Таким образом, та же степень повышения надежности, которая требует увеличения числа элементов штриха Шеффера в 50 000 — 70 000 раз, достигается увеличением числа ненадежных реле только в 80—120 раз.

Если теперь применить нижние оценки числа контактов, даваемые теоремой 7, то оказывается, что возрастание числа контактов не может быть сильно уменьшено. В самом деле, для получения конечной вероятности $8 \cdot 10^{-18}$ необходимо не менее 55 контактов (тогда как было использовано 81). Для получения $2 \cdot 10^{-19}$ необходимо не менее 66 (использовано 100), для $2,2 \cdot 10^{-21}$ — не менее 81 (использовано 121). Во всех этих случаях нижняя оценка очень близка к двум третям числа фактически использованных контактов.

Интересно разобраться в том, почему применение реле требует такого малого увеличения числа элементов по сравнению с тем, что представляется необходимым в случае элементов штриха Шеффера. Конечно, не доказано¹⁾, что способ повышения надежности, используемый фон Нейманом, приводит к наилучшему использованию этих элементов. Одно различие между этими двумя типами элементов, которое может объяснить получающееся расхождение, состоит в следующем. При синтезе устройств на обоих типах элементов осуществляется два процесса: во-первых, дублирование переменных, т. е. пересылка одной и той же переменной из одной части устройства в другие, и, во-вторых, формирование логических комбинаций и функций нескольких переменных. В случае ненадежных реле ошибки возникают при дублировании переменных. Состояния различных контактов реле не вполне соответствуют состоянию обмотки, а подвержены случайным отклонениям. Однако логические комбинации образуются без ошибок; последовательная комбинация абсолютно надежно реализуется схемой «и», параллельная комбинация — схемой «или».

В случае использования нейроноподобных элементов штриха Шеффера положение обратно. Можно любое число раз сдублировать данную переменную простым разветвлением линии, отходящей от соответствующего элемента. Все они предполагаются идентичными. Однако когда формируются логические комбинации двух или более переменных в элементе штриха Шеффера или в смесителе, появляются случайные ошибки. Далее, в обоих типах устройств статистически ненадежные части должны контролироваться логическими операциями, включающими в себя своего рода голо-сование.

«Sed quis custodiet ipsos custodies?»²⁾ В случае реле «сторожа» вне подозрений; в случае нейроноподобных элементов «сторожа» должны бдительно контролироваться другими «сторожами». Это принципиальное различие может объяснить разницу в увеличении числа элементов и, возможно, также тот факт, что в случае использования элементов штриха Шеффера для их контроля необходима некоторая степень надежности. В случае использования реле этого не требуется.

Интересно в этой связи, что для очень большого повышения надежности схемы фон Неймана требуют меньшего числа элементов, чем описанные нами. Если взять, как и раньше, $\delta_1 = \delta_2$ и $a = 1 - c = 0,005$, то, по фон Нейману, требуется приблизительно

1) Не доказано также, что методы, используемые в данной статье, дают наилучший возможный способ повышения надежности в случае релейных схем. Возможно, что более эффективным является дублирование целых схем, а не замена каждого реле отдельной схемой.

2) «Но кто сторожит самих сторожей?» — *Прим. перев.*

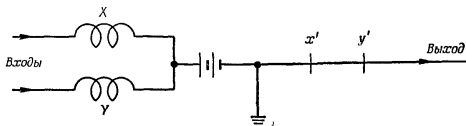
(беря логарифм его равенства (27), полагая $p = a$ и отбрасывая член более высокого порядка $\log \frac{6,4}{\sqrt{n}}$)

$$n \approx -3500 \log \delta_1.$$

Для наших схем имеем

$$n \approx \frac{1}{4} (\log \delta_1)^2.$$

Эти кривые пересекаются при $\delta_1 \approx 10^{-14.000}$! Если требуется большое повышение надежности, то восстанавливающие схемы будут требовать меньшего числа элементов.



Р и с. 19. Метод реализации штриха Шеффера с помощью двух реле.

Однако эффект элемента штриха Шеффера может быть получен путем использования схемы рис. 19, в которой два реле X и Y , каждое из которых имеет по одному размыкающему контакту, обозначенному соответственно x' и y' , объединены в такой элемент. Такая схема в качестве элемента штриха Шеффера имеет надежность, лишь немного меньшую, чем ее компоненты, и она может быть использована в схемах, описанных фон Нейманом. Таким образом, для получения весьма высокой надежности, когда схемы фон Неймана имеют преимущество, можно строить схемы с увеличением числа элементов порядка первой степени логарифма δ_1 , причем это увеличение будет немного больше, чем в случае элементов штриха Шеффера.

Обратный метод, т. е. попытка использовать корректирование релейного типа для элементов штриха Шеффера, оказывается неприменимым. По-видимому, из таких элементов никаким способом невозможно построить элемент, который работал бы совершенно так же, как реле.

Реле с неопределенным временем срабатывания

Функцию вероятности $h(p)$, связанную с контактной схемой, можно интерпретировать и некоторым другим образом, отличным от приведенного выше. Предположим, что каждый из контактов в схеме принадлежит своему реле и что эти реле имеют неопределен-

ное время срабатывания. Пусть $\varphi(t)$ — интегральная функция распределения для одного из этих реле; если реле возбуждено при $t=0$, то $\varphi(t)$ — вероятность того, что контакт будет замкнут в момент t . Одна и та же функция $\varphi(t)$ предполагается применимой ко всем реле, и последние предполагаются статистически независимыми. Тогда интегральное распределение для двухполюсной схемы имеет вид $h[\varphi(t)]$. Объяснение этого факта состоит в том, что для каждого момента времени t_1 каждый из контактов может считаться контактом ненадежного реле с вероятностью замыкания $p = \varphi(t_1)$. Следовательно, вероятность схемы быть замкнутой в момент t_1 равна $h(p) = h[\varphi(t_1)]$.

Это утверждение справедливо и в том случае, когда реле одновременно имеют неопределенное время срабатывания и являются ненадежными. В этом случае $\varphi(t)$ изменяется от 0 до одного из чисел c и a (см. рис. 1), когда t изменяется от 0 до $+\infty$.

Полученный нами результат о производной функции вероятности $h(p)$ показывает в этой интерпретации, что реле с неопределенным временем срабатывания могут быть использованы для синтеза схем с точным временем срабатывания. Тот факт, что график функции $h(p)$ пересекает диагональ самое большее один раз, показывает, что в некотором смысле время срабатывания любой двухполюсной схемы менее неопределенно, чем время срабатывания ее компонент. Таким образом, эта интерпретация показывает, что замена отдельных контактов такими двухполюсными схемами ведет к улучшению временных границ и снижению вероятности ошибок, связанных с очередностью срабатывания реле.

Синтез надежных схем в целом

До сих пор решалась задача синтеза схем, которые должны работать как простой надежный контакт. Однако не очевидно, что если заменить контакты большой схемы надежными схемами, то полученная большая схема будет вести себя обязательно надежно. Трудность состоит в том, что замещение отдельных контактов надежными схемами вводит возможность некоторых условий неодновременности срабатывания реле, которая, возможно, может явиться причиной таких ошибок, которые отсутствовали в первоначальной схеме. Например, одна из наших надежных схем может размыкаться и замыкаться в течение различного времени при переходе ее реле из возбужденного состояния в невозбужденное. Если такая схема является счетчиком импульсов, следующих друг за другом через небольшие промежутки времени, то очевидно, что могут возникнуть ошибки. Именно поэтому невозможно во всех случаях оправдать использование таких надежных схем. Имеется, однако, много типов схем, в которых этот эффект не встречается, и есть

основание полагать, что даже тогда, когда этот эффект имеет место, он фактически не страшен. Приведем несколько примеров таких случаев и их обоснование.

1. Схемы направленного действия. Под такими схемами подразумеваются релейные схемы, в которых реле могут быть распределены по уровням. Реле в уровне n управляются контактами входных реле и контактами реле низших уровней. Такой тип построения схемы ведет к тому, что в ней отсутствует память или обратная связь и что временные условия срабатывания реле несущественны для их окончательных состояний. Легко видеть, что если такая схема содержит N контактов и каждый из них заменяется надежной схемой с вероятностью неправильной работы менее P , то схема в целом будет иметь вероятность ошибки, меньшую чем NP . Эта верхняя оценка получается просто сложением вероятностей неисправности отдельных контактов и поэтому неудовлетворительна.

2. Синхронные релейные схемы. Под такими схемами понимаются схемы, срабатывающие в дискретные моменты времени и действующие подобно нейронным моделям Мак-Каллока и Питтса¹⁾ или подобно «селекторам» фирмы ИБМ. Точнее, контакты могут замыкаться или размыкаться только в моменты времени, кратные T , и если в момент nT реле возбуждено (или не возбуждено), то в момент $(n+1)T$ контакты имеют вероятность c (или a) быть замкнутыми. Схемы, построенные из таких элементов, даже если они должны иметь память и обратную связь, могут быть сделаны надежными посредством описанных приемов. Нетрудно показать, что вероятность ошибки в процессе работы такой схемы, если несколько расширить описанный ранее метод, меньше PND , где TD — время работы схемы, а P и N имеют тот же смысл, что и выше.

3. Описанным нами методом можно сделать надежными большое число обычных релейных схем, даже в случае наличия в них обратной связи и памяти. Как было показано, это происходит потому, что в надежных релейных схемах разброс времени срабатывания уменьшается. Если контактная схема становится все более и более надежной, то вместе с тем она и действует, подобно реле, с все более и более точно определенным временем срабатывания. Если происходит какое-нибудь ошибочное размыкание или замыкание, то оно с большой вероятностью ограничивается экстремально коротким временем, что мало влияет на точность работы схемы в целом. Таким образом, с большой вероятностью даже ошибочное срабатывание

¹⁾ McCulloch W., Pitts W., *Bull. Math. Biophysics*, 5 (1942), 115 (русский перевод в сб. «Автоматы», ИЛ, М., 1956).

отдельных элементов не приводит к нарушению работы всего устройства; с другой стороны, применение описанных методов дает значительное повышение надежности и уменьшение разброса времени срабатывания отдельных реле.

Приложение

Верхние оценки для вероятностей ошибок в случае гамакообразных l - w -схем

В нашем основном методе синтеза для получения требуемой надежности использовалась гамакообразная схема три-на-три и ее итерация с собой несколько раз, если это было необходимо. На основании этого метода были легко получены верхние оценки вероятностей ошибок. Возможно, однако, что более эффективным является применение итераций гамакообразных схем больших размеров.

Покажем, что если значения a , c , δ_1 , δ_2 заданы в прежнем смысле, то для гамакообразной схемы длины l и ширины w имеем неравенства

$$\delta_1 \leq \left(\frac{1 - \sqrt{1 - 16a^2}}{4a} \right)^{l-1} wa,$$

$$\delta_2 \leq \left(\frac{1 - \sqrt{1 - 16(1-c)^2}}{4(1-c)} \right)^{w-1} l(1-c).$$

При малых a и $1 - c$ правые части этих неравенств приблизительно равны $(2a)l^{-1}wa$ и $[2(1-c)]^{w-1}l(1-c)$.

Для доказательства этих неравенств рассмотрим сначала бесконечную схему, изображенную на рис. 20. Все наклонные отрезки изображают на ней контакты, для каждого из которых вероятность быть замкнутыми равна p . Схема предполагается неограниченно продолженной вверх, вниз и вправо. Вертикальная линия L слева изображает замыкание накоротко. Пусть P_1 — вероятность того, что один из узлов на расстоянии 1 от L будет соединенным с линией L . Аналогично P_2 — такая же вероятность для расстояния 2 и P_n — для расстояния n . Требуется найти верхнюю оценку для P_n .

На рис. 21 узел a , находящийся на расстоянии n , будет соединен с L , если узел b соединен с L и контакт c_1 замкнут; или если

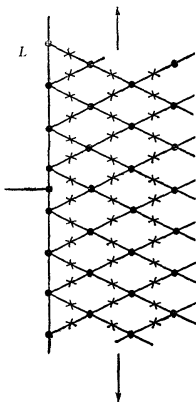
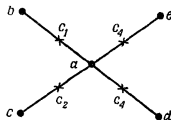


Рис. 20. Бесконечная гамакообразная схема.

узел c соединен с L и контакт c_2 замкнут; или если узел d соединен с L и контакт c_3 замкнут; или если узел e соединен с L и контакт c_4 замкнут. Верхняя оценка вероятности первого из этих четырех случаев есть pP_{n-1} . То же самое верно для второго случая.



Р и с. 21. Один узел бесконечной гамакообразной схемы.

Последние два случая имеют верхние оценки pP_{n+1} . Если их сложить, наша оценка вероятности P_n только увеличится (так как одновременно могут иметь место оба эти случая). Поэтому можно написать

$$P_n < 2pP_{n-1} + 2pP_{n+1}. \quad (21)$$

Очевидно, что $P_{n+1} \leq P_n$. Используя это, получаем из предыдущего неравенства

$$\begin{aligned} P_n &< 2pP_{n-1} + 2pP_n, \\ P_n &< \frac{2p}{1-2p} P_{n-1}. \end{aligned} \quad (22)$$

Это соотношение, если вспомнить, что $P_0 = 1$, дает

$$P_n < \left(\frac{2p}{1-2p} \right)^n.$$

Можно, однако, пользуясь любопытным итеративным методом, получить лучшую оценку для P_{n+1} . Именно из неравенства (22) получается $P_{n+1} \leq \frac{2p}{1-2p} P_n$. Подставляя эту оценку для P_{n+1} в неравенство (21), имеем

$$\begin{aligned} P_n &< 2pP_{n-1} + \frac{4p^2}{1-2p} P_n, \\ P_n &< \frac{2p}{1 - \frac{4p^2}{1-2p}} P_{n-1} = \frac{2p(1-2p)}{1-2p-4p^2} P_{n-1}. \end{aligned}$$

Это неравенство снова можно использовать для улучшения оценки. Чтобы найти результат бесконечного приближения в этом процессе, допустим, что на j -м шаге имеем соотношение

$$P_{n+1} < a_j P_n.$$

Подставляя его в неравенство (21), получаем

$$P_n < 2pP_{n-1} + 2p\alpha_j P_n$$

и, следовательно,

$$P_n < \frac{2p}{1-2p\alpha_j} P_{n-1},$$

$$\alpha_{j+1} = \frac{2p}{1-2p\alpha_j}.$$

Эта гиперболическая зависимость между α_{j+1} и α_j представлена на рис. 22. Если $p < 1/4$, то гипербола пересекает прямую $\alpha_{j+1} = \alpha_j$

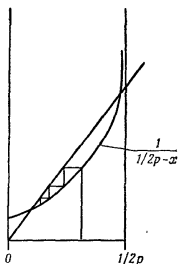


Рис. 22. Процесс итерации, используемый для улучшения оценки.

в двух точках, абсциссы которых являются корнями уравнения

$$2p\alpha^2 - \alpha + 2p = 0,$$

$$\alpha = \frac{1 \pm \sqrt{1-16p^2}}{4p}.$$

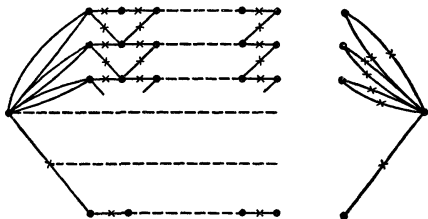
Нижняя точка пересечения, соответствующая корню $\frac{1 - \sqrt{1-16p^2}}{4p}$, является неподвижной точкой этого итеративного процесса. Начиная с любого α_0 , лежащего между нулем и единицей, α_n будет стремиться в пределе к этому корню «ступенчатым образом», как показано на рис. 22. Отсюда следует, что

$$P_n \leq \left(\frac{1 - \sqrt{1-16p^2}}{4p} \right) P_{n-1}$$

и, следовательно,

$$P_n \leq \left(\frac{1 - \sqrt{1 - 16p^2}}{4p} \right)^n. \quad (23)$$

Рассмотрим теперь гамакообразную схему с размерами l и w . Ее можно мысленно разбить на две части, как это изображено



Р и с. 23. Гамакообразная схема с размерами l и w , разделенная на две части.

на рис. 23. Вероятность замыкания такой схемы оценивается неравенством

$$h(p) \leq P^{(1)}(2p - p^2) + P^{(2)}(2p - p^2) + \dots$$

Последний член имеет вид $P^{(s)}(2p - p^2)$ или $P^{(s)}p^{(s)}$ в зависимости от того, четно w или нечетно. В этом выражении $P^{(j)}$ есть вероятность того, что левая половина схемы замкнута до j -го узла сверху в самом правом ряду. Множители $2p - p^2$ являются вероятностями замыкания двух параллельно соединенных контактов. Члены суммы поэтому соответствуют различным способам замыкания схемы. Если их сложить, то получится верхняя оценка вероятности.

Ясно, что каждое $P^{(i)} \leq P_{i-1}$, так как можно получить схему рис. 20 из левой части рис. 23 путем добавления контактов, которые, конечно, только увеличивают вероятность замыкания. Следовательно, из неравенства (23) вытекает, что

$$P^{(i)} < \left(\frac{1 - \sqrt{1 - 16p^2}}{4p} \right)^{i-1}.$$

Используя это неравенство, а также усиливая его отбрасыванием отрицательных членов с p^2 , приходим к искомому результату

$$\delta_1 < \left(\frac{1 - \sqrt{1 - 16a^2}}{4a} \right)^{l-1} wa.$$

Двойственной к гамакообразной схеме с размерами l и w является гамакообразная схема с размерами w и l . В силу двойственности получаем

$$\delta_2 < \left(\frac{1 - \sqrt{1 - 16(1-c)^2}}{4(1-c)} \right)^{w-1} l(1-c).$$

Можно заметить, что для малых a и $1-c$ эти верхние оценки становятся приблизительно равны $(w/2)(2a)^l$ и $(l/2)(2(1-c))^w$. Мы предполагаем, но не можем доказать, что для всех a и c

$$\delta_1 < \frac{w}{2} (2a)^l,$$

$$\delta_2 < \frac{l}{2} (2(1-c))^w.$$

ИСПОЛЬЗОВАНИЕ МАШИНЫ ДЛЯ ПРОЕКТИРОВАНИЯ ПЕРЕКЛЮЧАТЕЛЬНЫХ СХЕМ¹⁾

Введение

Можно очень просто описать некоторые операции, помогающие при проектировании релейных переключательных схем или схем других типов, и сконструировать машины, которые будут выполнять эти операции быстрее и точнее, чем человек. Представляется, что машины такого типа окажут пользу тем, кто работает над проектированием подобных схем.

Рассматриваемая машина, носящая название анализатора контактных схем, предназначается для использования в связи с проектированием двухполюсных схем, состоящих из контактов не более чем на четырех реле. Принципы, на которых основана эта машина, не ограничены двухполюсными схемами или четырьмя реле, хотя машина с более широкими возможностями будет работать дольше. Прибавление каждого нового реле к рассматриваемым схемам почти удвоит размеры машины и учетверит время ее работы. Этот тип машин неприменим к многотактным (последовательностным) схемам; поэтому он полезен только для работы с частями таких схем, содержащими контакты и не содержащими обмоток реле.

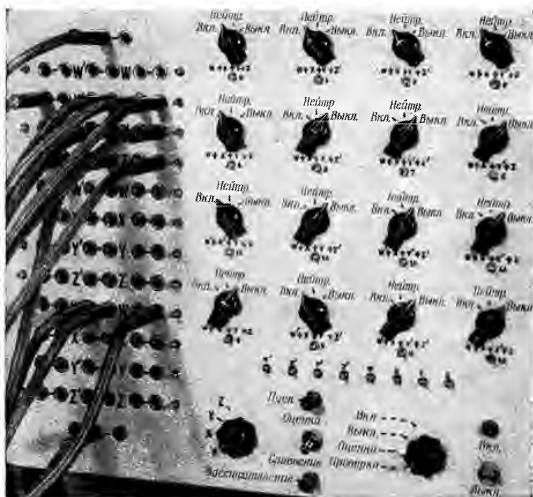
Работа машины

Машина, как показано на рис. 1, имеет шестнадцать трехпозиционных переключателей. Эти переключатели используются для задания условий функционирования схемы. Каждый переключатель соответствует одному из $2^4 = 16$ состояний, в которых могут находиться четыре реле. Под переключателем номер два в верхнем правом углу, например, написано $w + x + y' + z$, что соответствует такому состоянию схемы, при котором реле w , x и z возбуждены, а реле y не возбуждено.

Три положения переключателя соответствуют требованиям, которые можно предъявить к схеме при соответствующем состоянии

¹⁾ Shannon C., Moore E., Machine aid for switching circuit design, *Proceedings of the IRE*, 41, № 10 (1953), 1348.

реле. Поскольку любая одноконтактная схема принимает только два значения (разомкнуто или замкнуто), включение третьего



Р и с. 1 Внешний вид передней панели анализатора.

значения (безразличного, незначащего, или пустого, как его называют¹⁾) заслуживает объяснения. Если машина, частью которой должна быть проектируемая контактная схема, позволяет этим реле принимать только часть 2^н комбинаций, в которых могут находиться n реле, то не имеет значения, что делает схема во всех остальных случаях. Оператор указывает это обстоятельство, ставя

¹⁾ В нашей литературе принято название «неиспользуемого». — Прим. ред.

соответствующие переключатели в «безразличное» положение. При такой формулировке требований имеется большой выбор схем, удовлетворяющих им, и поэтому больше вероятности того, что окажется достаточной схема с меньшим числом контактов, чем в полностью определенном случае. Таким образом, шестнадцать трехпозиционных переключателей позволяют оператору не только требовать построения схемы, реализующей определенную функцию сопротивления, но и дают машине большую свободу в тех случаях, когда функция не определена полностью.

Для того чтобы машина этого типа могла оперировать с n реле (рассматриваемая машина была сделана для $n = 4$), требуется 2^n таких переключателей, соответствующих 2^n состояниям, которые могут принимать n реле. В каждом из этих состояний схема может быть либо замкнута, либо разомкнута, так что имеется 2^{2^n} функционально различных схем. Но поскольку каждый переключатель имеет 3 положения, то имеются 3^{2^n} различных требований, определяемых при помощи переключателей, что в случае $n = 4$ составляет 43 046 721. Поэтому число задач, с которыми может иметь дело анализатор, достаточно велико даже в случае только четырех реле.

Левая половина переднего щита машины представляет собой коммутационную доску, на которой можно представить анализируемую схему (см. рис. 1). У каждого из четырех реле, w , x , y и z , имеются 3 переключающих контакта, соединенных с гнездами на переднем щите, а гнезда, представляющие полюсы схемы, находятся сверху и внизу. Используя соединительные шнуры, можно задать на коммутационной доске любую схему, содержащую не более трех переключающих контактов от каждого из четырех реле. Это число контактов достаточно для получения схемы, реализующей любую переключательную функцию от четырех переменных.

Если условия для схемы были заданы на шестнадцати переключателях и схема была набрана на коммутационной доске, то релейный анализатор готов к действию.

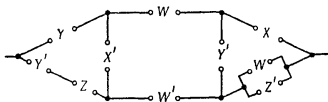
Когда главный переключатель и переключатель оценки-сравнения находятся в положении «оценка», нажатие пусковой кнопки заставит анализатор оценить заданную схему, т. е. указать, в каких состояниях схема замкнута. Это делается при помощи включения соответствующих индикаторных лампочек.

Когда переключатель оценки-сравнения поставлен в положение «сравнить», анализатор проверяет, не расходится ли схема с условиями, заданными на переключателях. Расхождение указывается индикаторной лампочкой, соответствующей данному состоянию. Если переключатель находится в состоянии «замкнуто», а схема при этом разомкнута, или наоборот, то это указывается соответствующими лампочками. Однако если переключатель находится в нейтраль-

ном положении, то такое расхождение не отмечается независимо от состояния схемы. Хотя положение «сравнить» и дает информацию, равноценную информации положения «оценить», однако оно является более удобным для обнаружения ошибок.

После того как найдена схема, полностью отвечающая заданным требованиям, главный переключатель ставится в положение «испытание на замыкание» и вновь нажимается пусковая кнопка. Теперь машина определяет, нельзя ли накоротко замкнуть некоторые контакты схемы так, чтобы схема при этом отвечала поставленным требованиям. Возле контактов, которые оказались лишними, загораются лампочки.

Быть может, читателю покажется странным, что вообще нужна помощь машины для отыскания контакта, который можно замкнуть, не влияя при этом на свойства схемы. Хотя это справедливо для простых случаев, в более сложных схемах такие лишние элементы часто далеко не очевидны (в частности, так бывает при наличии ряда нейтральных положений переключателей, поскольку упрощенная схема может быть функционально отличной от первоначальной, отличаясь от нее только в состояниях, соответствующих нейтральным положениям переключателей). В этих случаях часто очень трудно найти способ упрощения.



Р и с. 2. Анализатор упростил эту схему, выкинув один контакт, за две минуты. Можете ли вы сделать это так же быстро?

Анализатор помогает также в случае, когда анализируемая схема представляет собою мостик, так как бывает трудно проследить в нем все пути. Схема, изображенная на рис. 2, является примером такой схемы, которую не удавалось эффективно синтезировать до тех пор, пока ее не проверили на анализаторе. Анализатор определил менее чем за две минуты (считая время, необходимое для задания схемы на коммутационной доске), что в этой схеме можно замкнуть один из контактов. Может ли человек решить эту задачу в столь же короткий срок?

После того как произведено испытание на замыкание, главный переключатель ставится в положение «испытание на размыкание», что позволяет анализатору произвести другое аналогичное испытание, на этот раз поочередно размыкая контакты.

Выбор именно этих двух типов изменения схемы был сделан вследствие легкости их осуществления, а также потому, что в случае удачи каждый из них сокращает число контактов в схеме. Есть и другие пути упрощения, включая различные преобразования схемы, которые было бы желательным поручить машине. Они потребовали бы большего времени и более сложного оборудования, но, быть может, позволили бы чаще достигать успеха при машинном решении подобных задач. Используя описанные методы, можно было бы построить машину, способную эффективно синтезировать схемы, исходя из основных принципов, быть может, начиная с полного разложения¹⁾ для реализуемой функции, а затем сокращая его шаг за шагом.

Такая машина действовала бы довольно медленно (даже если бы ее построить так, чтобы она работала с электронными скоростями, и то она работала бы медленно). Не разрабатывался еще в полной мере вопрос о том, осуществима ли подобная машина практически. Однако тот факт, что такая машина теоретически возможна, несомненно, представляет интерес, независимо от того, построит ее кто-нибудь или нет.

Другой вопрос, представляющий теоретический интерес, состоит в следующем. Можно ли построить логическую машину, которая могла бы спроектировать улучшенный вариант самой себя, или, быть может, построить некую машину, общая цель которой была бы сложнее, чем ее собственная? Представляется, что в подобной машине нет логического противоречия, хотя, прежде чем с уверенностью заниматься подобным проектом, необходимы большие достижения в общей теории автоматов.

Но вернемся к анализатору релейных схем. Конечная его операция производится в положении главного переключателя «доказать». Когда нажимается пусковая кнопка и другой переключатель последовательно передвигается через четыре положения w , x , y и z , загораются некоторые из восьми ламп w , w' , x , x' , y , y' , z , z' . Анализатор устанавливает, какие типы контактов необходимы для реализации функции, используя метод приведения к функциям одного переменного, который будет объяснен в будущих работах²⁾. Анализатор здесь не принимает во внимание, какая схема была задана на коммутационной доске, и рассматривает только функцию,

¹⁾ См. стр. 16.—Прим. ред.

²⁾ По-видимому, авторы имеют в виду следующее. Пусть $f(x_1, \dots, x_n)$ — некоторая булева функция и x_i — один из ее аргументов. Если среди 2^{n-1} функций $f(\sigma_1, \dots, \sigma_{i-1}, x_i, \sigma_{i+1}, \dots, \sigma_n)$ встретится функция x_i , то (всякая) схема для f должна иметь замыкающий контакт реле x_i , а если встретится функция $\bar{x}_i$, то (всякая) схема для f должна иметь размыкающий контакт реле x_i . — Прим. ред.

определяемую положениями шестнадцати трехпозиционных переключателей. Если каждая схема, удовлетворяющая заданным условиям, требует размыкающего контакта реле w , то загорится лампочка w' и т. д.

Если, например, загорелись семь из восьми лампочек, то любая схема для функции требует по меньшей мере семи контактов, и если действительно имеется схема с семью контактами, то машина полностью доказала, что схема минимальна. Схемы, для которых машина может дать подобное полное доказательство, довольно обычны, хотя имеются также схемы (минимальность которых может быть доказана более тонкими методами), доказать минимальность которых машина не могла. Примером является схема, приведенная на рис. 2. Эта схема может быть упрощена анализатором до девяти контактов, но в положении «доказать» анализатор только указывает, что необходимо по меньшей мере восемь контактов. Другими методами можно показать, что схема с девятью контактами является минимальной. Но во всяком случае анализатор всегда дает математически точную нижнюю оценку необходимого числа контактов.

Небольшие размеры и портативность машины определяются тем, что в ней были использованы как реле, так и элементы электронных схем. Неоновые лампочки особенно удобны там, где требуется небольшой элемент памяти в совокупности с устройством визуальной сигнализации, а реле и селекторные переключатели — там, где требуется обеспечение очередности срабатывания элементов и наличие внутренних связей между ними при малом весе и габаритах деталей. В целом анализатор релейных схем имеет в качестве логических элементов только двадцать четыре реле, два селекторных переключателя, сорок восемь миниатюрных неоновых лампочек и четырнадцать полупроводниковых диодов.

Для тех, кто знаком с универсальными вычислительными машинами, может представить интерес сравнение этого метода решения данной задачи на такой маленькой специализированной машине с более удобным методом программирования этой же задачи для решения на быстродействующей универсальной вычислительной машине. Одно из основных различий между этими двумя методами состоит в том, что анализируемые схемы задаются на анализаторе непосредственно. На универсальной вычислительной машине было бы необходимо иметь символическое описание схемы, быть может, в форме числового кода, описывающего соединения схемы и обозначающего типы контактов в различных частях схемы при помощи ряда чисел, находящихся в памяти машины. Анализатор релейных схем представляет схему более непосредственным и естественным путем, так как имеется копия схемы на коммутационной доске на передней панели машины.

Разница в степени непосредственности представления схемы в машине имеет два следствия. Во-первых, обычную вычислительную машину использовать несколько труднее, так как перевод соединений схемы в код и ввод его в машину был бы сложнее и дольше, чем непосредственное задание схемы на коммутационной доске. Во-вторых, имеется разница в числе логических операций (и отсюда, косвенно, во времени), необходимых для этих двух машин. Для выполнения основного этапа работы по определению того, замкнута или разомкнута схема (или ее модификация, полученная размыканием или замыканием одного из контактов) при некотором состоянии реле, необходимо время работы только одного реле анализатора контактных схем, а для выполнения этого этапа на универсальной вычислительной машине потребовалось бы много раз повторять различные вспомогательные операции. Хотя имеются различные способы программирования задачи, в типичном случае машина должна сначала с помощью подпрограммы выяснить, замкнут или разомкнут данный контакт, и повторить эту операцию для каждого контакта схемы, а затем с помощью другой подпрограммы анализировать каждую вершину схемы. В целом это, пожалуй, потребовало бы выполнения нескольких сотен команд, хотя при достаточно умелом программировании число команд можно было бы сократить до 100. Поскольку каждая команда требует для выполнения приблизительно в 100 раз больше времени, чем единичная логическая операция (т. е. время одного такта, если машина действует по принципу часового механизма), то, что требует отрезка времени одной операции на первой машине, требует около 10 000 отрезков на другой.

Поскольку 10 000 приблизительно равно отношению между скоростями срабатывания реле и электронной лампы при выполнении логических операций, этот выигрыш около 10 000 в результате непосредственного представления схемы анализатором позволяет релейной машине решать задачи такого типа столь же быстро, как и электронная вычислительная машина.

Большая разница в скоростях этих двух машин не типична для всех задач, поскольку обычная задача численного анализа решается раз в десять быстрее на специализированной машине (так как операции умножения и деления составляют около одной десятой общего времени решения задачи). Однако в комбинаторных задачах представляется возможным получить огромный выигрыш во времени при использовании специализированных, а не универсальных вычислительных машин. Это значит, что универсальные машины в действительности не являются универсальными, а специализируются так, чтобы в основном решать задачи анализа. Верно, конечно, что так называемые универсальные машины способны решать такие комбинаторные задачи, но их эффективность при подобном

использовании определенно низка. Проблемы, связанные с проектированием универсальных машин, пригодных для широкого круга комбинаторных задач, представляются очень трудными, хотя, несомненно, очень интересными теоретически.

Заключение

Интересной особенностью анализатора контактных схем является его способность непосредственно оперировать с логическими схемами в терминах трехзначной логики. Значительный интерес представили бы методы для простого действия с подобной логикой на бумаге, поскольку ее можно непосредственно использовать при синтезе экономичных переключательных схем. Хотя подобные методы еще не разработаны, такие машины, как описанная здесь, могут иметь значение в связи с проблемами трехзначной логики.

Независимо от того, будет полезен этот конкретный тип машин при проектировании реальных контактных схем или нет, возможность создания машин, помогающих при логическом синтезе, может способствовать облегчению работы по синтезу переключательных схем. Так же, как логарифмическая линейка и современные типы цифровых вычислительных машин могут помочь в части работы, связанной с синтезом линейных электрических схем, такие машины, как описанная, могут со временем намного облегчить работу по синтезу логических схем.

Введение

В 1871 г. Самуэлем Батлером была закончена рукопись книги «Эревол», являвшейся одной из редких социальных сатир того времени. Три главы «Эревол», вначале появившиеся под названием «Дарвин среди машин», были пародией на книгу Дарвина «Происхождение видов». В сатирическом стиле описывает Батлер эволюцию машин. Он классифицирует машины по семействам, родам и видам, по способам питания, рассматривает их рудиментарные органы чувств, их механизмы воспроизведения и эволюционирования (если машина не эффективна, человек вынужден проектировать более эффективные), тенденции к атавизму, отмирание рудиментарных органов и даже проблему свободной воли у машин.

Перечитывая «Эревол» в настоящее время, находишь «Книгу машин» поистине пророческой. Действительно, все существующие и проектируемые сейчас вычислительные машины и системы управления все в большей степени приобретают способности людей и животных, выполняют их функции и фактически им эти новые качества присущи в гораздо большей степени, чем это в свое время предвидел Батлер.

Универсальные вычислительные машины создавались в основном для решения задач, связанных с числовыми расчетами. Поэтому для большинства из нас наиболее интересные возможности вычислительных машин и управляющих систем связаны не с осуществлением вычислительных операций, а с решением логических задач, переводом с одного языка на другой, проектированием схем, играми, управлением датчиками и манипуляторами и вообще выполнением сложных функций, свойственных нашему мозгу.

Неарифметическое применение вычислительных машин нельзя рассматривать просто как частный случай обычных арифметических вычислений. Как раз наоборот. Сто лет назад Чарльз Бэббидж сконструировал свою замечательную аналитическую машину, управлявшуюся с помощью перфокарт, подобно тому, как управлялись

¹⁾ Shannon C., Computers and automata, *Proc. IRE*, 41, № 10 1953), 1234.

шелкоткацкие станки «Жакард», — машину, которая просуществовала полстолетия¹⁾).

Наиболее крупной и наиболее надежной системой, предназначенной для обработки информации, является автоматическая телефонная станция. Наши заводы и фабрики заполнены «умными» машинами всех образцов и форм, осуществляющими почти невероятные операции. На железных дорогах и электростанциях установлены совершенные системы контроля, предупреждения несчастных случаев и ошибок, допускаемых человеком.

Однако все это — автоматы, предназначенные для специальных целей. Новым в применении машин для выполнения неарифметических операций является мысль о создании программно-управляемой вычислительной машины общего назначения (универсальной) — устройства, способного выполнять длинную последовательность элементарных операций, аналогичных выполняемым счетными машинами. Однако здесь элементарные операции относятся не только к числам, но и к физическим явлениям, операциям со словами, уравнениям, данным, полученным от различных датчиков, и почти к любым физическим и логическим величинам.

В данной статье дается краткое описание некоторых исследований в области неарифметического применения вычислительных машин и рассматриваются отдельные проблемы, возникающие в связи с этими исследованиями. Этот вопрос очень сложен, и в небольшой статье можно упомянуть лишь о некоторых результатах.

Мозг и вычислительные машины

Мозг часто сравнивали, иногда слишком восторженно, с вычислительными машинами. Мозг содержит около 10^{10} активных элементов, называемых нейронами. Так как передача нервных возбуждений осуществляется по принципу «все или ничего», нейроны имеют некоторое функциональное сходство с элементами двоичной вычислительной машины — реле, лампами или транзисторами. Правда, количество клеток в миллион раз (на 6 порядков) превышает количество элементов, используемых в самых сложных вычислительных устройствах. Мак-Каллок образно выразился, что вычислительная машина, которая имела бы столько ламп, сколько нейронов имеет человеческий мозг, потребовала бы для своего размещения Эмпайр Стейт Билдинг, Ниагарский водопад для обеспечения ее энергией и Ниагару для охлаждения. Использование полупроводников в такой машине значительно уменьшило бы ее размеры, необходимую для нее энергию и количество воды для ее

¹⁾ См. B a b b a g e C., *Passages from the Life of a Philosophy*, London, 1864. — Прим. перев.

охлаждения; так, например, для нее потребовалась бы мощность порядка нескольких сотен киловатт (мозг потребляет около 25 вт), размеры ее уменьшились бы (при условии компактного монтажа) до размеров обычного дома. Можно также сказать, что увеличение быстродействия электронных элементов, скажем в 10^3 раз, могло бы частично уменьшить размеры такой вычислительной машины.

Сравнения такого рода можно считать весьма преувеличенными, ибо наши сведения о функциях головного мозга, несмотря на большую и плодотворную исследовательскую работу в этой области, до сих пор весьма примитивны. Так, например, до настоящего времени остается открытым вопрос о том, достаточны ли наши сведения о нейроне для точного анализа его функций. Случайная структура на нейронном уровне (в числовом отношении), расположения и связи между нейронами наводят на мысль, что на данной стадии важна только статистика и что, следовательно, прежде чем конструировать математическую модель, необходимо получить усредненные данные о локальной структуре и функционировании мозга.

Обычно подчеркивают сходство мозга и вычислительных машин. Тем не менее различие, существующее между ними, пожалуй, более существенно, так как оно позволяет сделать заключение о важных особенностях мозга, не свойственных нашим лучшим моделям.

Наиболее важными из этих особенностей являются следующие.

1. *Разница в размерах.* Превышение количества составных частей на шесть порядков до такой степени далеко от обычной практики, что делает экстраполяцию функций почти бессмысленной.

2. *Разница в структуре.* Совершенно очевидно, что локально беспорядочная структура нервной системы в корне отличается от точного монтажа искусственных автоматов, где какое-нибудь одно ошибочное соединение может вызвать неполадки в работе всего механизма. Мозг имеет такое строение, что работа его в целом не зависит от точности микроструктуры.

3. *Разница в надежности.* Мозг может функционировать надежно десятки дней, без серьезного нарушения в работе (сравните, например, с бессмысленной неразберихой, выдаваемой вычислительной машиной в случае сбоя), несмотря на то, что его элементы сами по себе, вероятно, не более надежны, чем те, которые используются в вычислительных машинах.

4. *Разница в логической структуре.* Разница здесь настолько велика, что ее трудно изложить. Мозг является в значительной мере самоорганизующейся системой. Он может достаточно хорошо приспосабливаться к большому числу различных условий. Мозг обладает замечательным свойством классификации, записи и выборки информации в памяти, а также способностью быстро осуществлять информационный поиск с помощью большого числа «координатных систем». Он может создавать устойчивые «сервосистемы», облег-

чающие реализацию сложных связей между сенсорным входом и моторным выходом. В противоположность этому наши цифровые вычислительные машины выглядят как ученые-схоласты. При вычислениях длинной цепи арифметических операций цифровые вычислительные машины очень значительно обгоняют человека. Когда же пытаются приспособить вычислительные машины для выполнения неарифметических операций, они оказываются неуклюжими и неприспособленными для такой работы.

5. *Разница во вводных и выводных устройствах.* Мозг обладает прекрасно сконструированными вводными устройствами — органами чувств, в частности ушами и глазами. Наши лучшие устройства, имитирующие эти свойства, такие, как читающий анализатор Шепарда для опознавания и выдачи печатных знаков и «Одри» — звукоанализирующее устройство, которое может распознавать звуки речи, представляя их при помощи десяти цифр, при сравнении с мозгом выглядят весьма жалко. На выходе мозг управляет сотнями мускулов и желез; руки имеют 60 степеней свободы. Сравните эту способность со способностью разработанного Массачусетским технологическим институтом фрезерного станка, который управляется цифровым устройством и имеет только три степени свободы. Большинство вычислительных машин не имеет почти никаких средств восприятия окружающей среды и воздействия на нее и работает в абстрактной обстановке чисел и операций над ними.

Машины Тьюринга

Основы математической теории цифровых вычислительных машин были разработаны А. М. Тьюрингом в 1936 г. в классическом труде «О вычислимых числах и их применении к проблеме разрешения». Он определил класс вычислительных машин, которые сейчас называются машинами Тьюринга¹⁾. Эти машины в основном состоят из бесконечной бумажной ленты и вычислительного устройства, которое имеет конечное число внутренних состояний и способно считывать с одной ячейки ленты, писать в этой ячейке и передвигать ленту на одну ячейку вправо или влево. В каждый данный момент времени вычислительное устройство находится в определенном состоянии и считывает то, что записано в данной ячейке ленты. Следующая операция определяется текущим состоянием и прочитанным символом. Эта операция состоит в принятии нового состояния и в написании нового символа (вместо уже прочитанного) или в смещении ленты вправо или влево. Машины этого типа могут производить операции над числами при наличии соответствующего кода для представления символов. Например, в машине Тьюринга

¹⁾ См. Клини С., Введение в метаматематику, ИЛ, М., 1959. Трахтенброт Б. А., Алгоритмы и машинное решение задач, Физматгиз, М., 1960. — Прим. перев.

печатаются окончательные ответы в виде двоичного кода в особых ячейках ленты при использовании других ячеек для промежуточных вычислений.

Можно показать, что такие устройства образуют исключительно большой класс вычислительных машин. Обычные цифровые вычислительные машины, которые не содержат случайных или вероятностных элементов, эквивалентны некоторой машине Тьюринга. Любое число, которое может быть вычислено с помощью этих машин или фактически получено любым обычным способом, может быть вычислено и с помощью соответствующей машины Тьюринга. Однако, как показал Тьюринг, существуют задачи, которые не могут быть решены, и некоторые числа, которые не могут быть вычислены ни на какой машине Тьюринга. Например, невозможно создать такую машину Тьюринга, которая могла бы в случае, если в нее ввести соответствующим образом закодированное описание другой машины этого же типа, всегда давать верный ответ на вопрос: может ли вторая машина продолжать бесконечно печатать символы в соответствующие места ленты, предназначенной для выдачи окончательного ответа. Первая машина может на каком-то этапе сорваться в бесконечный цикл. Существование таких механически неразрешимых проблем весьма интересно для логиков.

Тьюрингом была также высказана интересная мысль об универсальной машине, называемой сейчас универсальной машиной Тьюринга. Эта машина обладает следующим свойством: если на ленте этой машины напечатано соответствующим образом закодированное описание некоторой машины Тьюринга и первая машина пущена в соответствующей точке и в соответствующем состоянии, она действует как вторая машина, т. е. выполняет (обычно с гораздо меньшей скоростью) те же самые операции, которые выполняла бы моделируемая машина. Тьюринг показал, что такая универсальная машина может быть создана. Она может, конечно, вычислить любое вычислимое число. Большинство цифровых вычислительных машин при условии, что они имеют память неограниченной емкости определенного типа, эквивалентны универсальным машинам Тьюринга и могут в принципе моделировать любую другую вычислительную машину и вычислять любое вычислимое число.

Результаты работы Тьюринга были обобщены и переформулированы различным образом. Интересным обобщением является понятие *A*-вычислимости. Оно относится к классу таких машин Тьюринга, которые имеют ту дополнительную особенность, что они могут на какой-то стадии расчета задать вопрос второму устройству—«оракулу» и использовать его ответы для дальнейших вычислений. «Оракул» может дать ответы на некоторые вопросы, которые не в состоянии решать обычная машина Тьюринга, и, следовательно, его введение позволяет решать гораздо более широкий класс задач.

Логические машины

Булева алгебра может быть использована в качестве математического аппарата для изучения свойств релейных и переключательных схем. И, наоборот, можно решать задачи булевой алгебры и формальной логики с помощью простых релейных схем. Эта возможность была использована в ряде логических машин. Типичная машина этого вида, описанная Мак-Коллумом и Смитом, может оперировать с логическими формулами, содержащими до семи логических переменных. Необходимые отношения между этими переменными, определяемые данной логической задачей, вводятся в машину путем коммутации ряда типовых логических элементов.

В машине имеется шесть типов логических элементов, которые реализуют логические связи: «нет», «и», «или», «или же» (исключение или), «тогда и только тогда» и «если—то». Когда завершены все коммутации, машина проверяет все $2^7 = 128$ комбинаций значений основных переменных, останавливаясь на всех тех, которые удовлетворяют поставленным условиям. Машина также указывает количество «истинных» переменных в каждом из этих состояний. Мак-Коллум и Смит приводят следующую типовую задачу, которая может быть решена машиной.

Известно, что моряки всегда говорят правду, а инженеры всегда лгут. G и E — моряки. C заявил, что D — инженер. A сказал, что B утверждает, что C заявляет, что D говорит, что E настаивает, что F отрицает то, что G — моряк. Если A — инженер, то сколько инженеров среди упомянутых лиц?

Заслуживающей внимания особенностью этой машины является селективная система обратной связи, служащая для отыскивания частных решений логических уравнений без полного перебора всех возможных комбинаций. Это достигается путем введения элементов, которые устанавливают, удовлетворяются ли частичные логические отношения. Если эти отношения не удовлетворяются, то все входящие в них логические переменные вынуждены колебаться между их двумя возможными значениями. Таким образом, когда условия не удовлетворены, переменные, входящие в них, постоянно меняются, тогда как в случае, когда условия удовлетворены, переменные не меняются. Если все условия удовлетворены одновременно, машина останавливается, выдавая частное решение. Изменение переменных происходит только в том случае, когда условия, заданные соотношениями, не удовлетворены. Как правило, это приводит к решению быстрее, чем методический перебор всех случаев; но, как всегда бывает в том случае, когда вводится обратная связь, это ведет к возможности непрерывных колебаний. Мак-Коллум и Смит подчеркивают, что желательно процесс изменения перемен-

ных при небалансе в цепи обратной связи сделать возможно более случайным для того, чтобы дать возможность машине избежать «зацикливания»; это осуществляется с помощью системы реле, переключающихся случайным образом.

«Играющие машины»

Создание играющих машин является весьма интересной проблемой, привлекающей большое внимание. Правила игр дают весьма ограниченный круг условий, при которых машина может работать с совершенно определенной целью. Дискретная природа большинства игр хорошо согласуется с возможностями цифровой вычислительной техники, причем не требуется громоздкого перевода непрерывных величин в цифровую форму, необходимого для преобразования внешних физических воздействий в случае исполнительных и чувствующих автоматов.

Играющие машины могут быть разделены на следующие группы по признаку нарастания сложности.

1. *Машины типа «словаря».* Здесь ответный ход машины заранее определяется для каждой возможной позиции, которая может возникнуть в игре, и регистрируется в «словаре» или в функциональной таблице. Когда возникает определенная позиция, машина просто отыскивает соответствующий ход в словаре. Ввиду непомерных требований, предъявляемых к памяти, этот тип машин не представляет интереса и осуществим лишь для исключительно простых игр, как, например, «тик-так-тое»¹⁾.

2. *Машины, использующие строго определенные игровые формулы.* Для некоторых игр, например, таких, как Ним²⁾, известна строгая математическая теория; поэтому при помощи сравнительно простых соотношений можно в любой позиции, которая может создаться в игре, рассчитать ход, нужный для выигрыша партии. Механизация этих соотношений обеспечивает выбор наилучшего хода в игре.

3. *Машины, применяющие общие принципы приблизительной оценки.* Для большинства игр, представляющих интерес, неизвестно простое и точное решение, но существуют различные общие принципы игры, которые справедливы в большинстве позиций, воз-

¹⁾ Тик-так-тое (tic-tac-toe) — игра в крестики и нолики на квадратном поле из девяти клеток (3×3). Два игрока поочередно заполняют свободные клетки, один крестиками, другой ноликами. Цель игры — первым расположить три своих значка в один ряд (по строке, столбцу или диагонали). — Прим. перев.

²⁾ Ним — игра, заключающаяся в следующем: некоторое количество предметов (фишек, спичек) разбивается на несколько групп. Каждый игрок поочередно берет любое число предметов (не менее одного) из какой-нибудь группы, но каждый раз только из одной. Проигрывает тот игрок, который вынужден взять последний предмет. — Прим. перев.

никающих в процессе игры. Это относится к таким играм, как шашки, шахматы, бридж, покер и т. д. Могут быть сконструированы машины, которые используют эти общие принципы при оценке сложившейся ситуации. Поскольку нет непогрешимых принципов, то ни машины, ни люди не могут быть непогрешимы.

4. *Обучающиеся машины.* Здесь машине даются только правила игры и иногда элементарная тактика игры вместе с некоторыми методами улучшения этой тактики в результате опыта. Среди многочисленных методов, предложенных для включения в программу обучения, можно указать:

а) метод «испытаний и ошибок» с запоминанием удачных и устранением неудачных вариантов;

б) подражание более успешно играющему партнеру;

в) «обучение» путем поощрения или штрафа или путем сообщения машине сведений о характере допущенных ошибок;

г) анализ машиной ошибок, допущенных самой машиной, с целью выработки общих принципов поведения.

Было сконструировано много машин первых двух и лишь несколько машин третьего типа. Машина четвертого типа, обучающаяся игре, напоминает замечания Марка Твена о погоде. Это прямой вызов специалистам по программированию и по конструированию машин.

Из машин, относящихся к третьему типу (применяющих общие принципы), интерес представляют следующие две. Первая из них сконструирована Э. Ф. Муром и автором данной статьи для игры в настольную игру, известную под названием «Гекс». Эта игра происходит на доске, разделенной на правильные шестиугольники. Два игрока по очереди ставят черные и белые фигуры на незанятые шестиугольники. Вся доска в целом образует ромб¹⁾ и цель черных — соединить верхнюю и нижнюю стороны ромба непрерывной цепью черных фигур, а цель белых — соединить две боковые стороны ромба непрерывной цепью белых фигур. После изучения этой игры было установлено, что вполне хорошие ходы могут быть получены в результате следующего процесса. Создается поле двумерного потенциала, соответствующее игровой доске, с положительными зарядами вместо белых фигур и отрицательными зарядами вместо черных фигур. Верхняя и нижняя стороны являются отрицательными, а боковые стороны — положительными. Ход, который должен быть сделан, соответствует некоторой седловой точке этого поля.

Для того чтобы испытать эту тактику игры, был сконструирован прибор аналогового типа, состоящий из схемы на сопротивлениях и устройства для локализации седловой точки. Этот общий метод

¹⁾ В некоторых вариантах игры — шестиугольник. — *Прим. перев.*

с некоторыми улучшениями, подсказанными практикой, оказывается вполне состоятельным. Машина выигрывает 70% игр у своего партнера — человека. Она часто удивляла своих конструкторов, выбирая странные, на первый взгляд непонятные ходы, которые при дальнейшем анализе оказывались, однако, вполне оправданными. Обычно считают, что машины справляются, с длинными запутанными вычислениями и пасуют перед общей оценкой обстановки. Как ни странно, оценка позиций этой машиной была хорошей; основным ее недостатком является выбор заключительных ходов в комбинационной игре. Любопытно также отметить, что машина «игрок в Гекс» нарушает обычную процедуру счета, а именно в этой машине по существу дискретная задача решается с помощью аналоговой системы.

Недавно Стрэчи запрограммировал игру в шашки с помощью универсальной вычислительной машины, используя метод «общего принципа». Он использовал прием, предложенный автором этой статьи для игры в шахматы, — исследование возможных вариантов на несколько ходов и минимаксную оценку окончательных позиций. Ниже приводится простая партия, сыгранная машиной и Стрэчи с пояснительными замечаниями Стрэчи. (В скобках указаны позиции шашек противника, снимаемых при данном ходе.)¹⁾

В шахматной нотации партия запишется следующим образом:

Машина	Стрэчи	Машина	Стрэчи
1. a3—b4	b6—c5	18. c5—b6	c7—d6
2. b2—a3	f6—e5	19. b6—a7	d6—e5
3. g3—h4	h6—g5 ^{a)}	20. f2—g3 ⁱ⁾	e5—d4
4. h4 : f6 (g5)	e7 : g5 (f6)	21. g3—f4	d4—e3
5. e3—d4 ^{l b)}	c5 : e3 (d4)	22. a5—b6	e3—d2
6. d2 : h6 (g5; e3) ^{c)}	a7—b6	23. b6—c7 ^{j)}	b8 : d6 (c7)
7. c1—b2 ^{d)}	b6—c5	24. a7—b8D	d2—c1 D
8. f2—g3	c5—d4	25. b2—c3	c1—b2
9. g3—f4 ^{e)}	e5 : g3 (f4)	26. c3—b4	f6—g5? ^{h)}
10. h2 : f4 (g3)	d4—e3	27. a1 : e3 (b2)	g5 : e3 (f4)
11. e1—f2	e3—d2	28. c3—d4	e3—d2
12. b4—a5	d2—e1 D	29. b4—a5	d2—c1 D
13. f2—e3	e1—d2? ^{g)}	30. b8—a7 ⁿ⁾	c1—d2
14. g1—f2 ^{h)}	d2 : b4 (c3)	31. a7—b8 ^{m)}	d2—c3
15. a3 : e7 (d6; b4)	d8—f6	32. b8—c7 ⁿ⁾	c3 : e5 (d4)
16. f4—e5	f6 : d4 (e5)	33. a5—b6	f8—e7
17. e3 : c5 (d4)	g7—f6	34. c7—b8 ^{p)}	Партия окончена

¹⁾ Рассматриваемая игра в шашки отличается от общепринятой. Простые не могут ни ходить назад, ни брать ходом назад. Дамки приобретают право ходить назад, но только на одну клетку за ход.— *Прим. перев.*

Замечания Стрэчи:

- а) Пробный ход с моей стороны—единственная преднамеренная жертва шашки, которую я сделал. Я ошибочно думал, что этот ход вполне безопасен.
- б) Непредусмотренный мною.
- с) Лучше, чем f2 : f6 (e3; e5).
- д) Ход, сделанный наугад (нулевое значение). Показывает на отсутствие конструктивного плана.
- е) Другой ход, сделанный наугад, имеющий нулевое значение. Фактически гораздо лучше.
- ф) Плохой. В конечном счете позволяет мне пройти в дамки. Ход с3—d4 был бы лучше.
- г) Промах с моей стороны.
- h) Извлекающий прямую выгоду из моей ошибки.
- и) Плохой. Открывает путь дамке.
- j) Жертва для того, чтобы получить дамку (а не преградить мне путь в дамки). Хороший ход, но невозможный до того, как ход от b6 будет сделан наудачу.
- к) Другой неудачный ход с моей стороны.
- т) Бесцельный. Тактика ухудшается к концу игры.
- п) Слишком поздно.
- р) Бесполезный ход. Игра на этом заканчивается, поскольку исход ее очевиден.

Совершенно очевидно, что, не обладая качествами чемпиона, машина все же играет лучше многих людей. Стрэчи указывает на различные недостатки в программе, особенно в определенных позициях в конце игры, и предлагает возможные улучшения.

Обучаемые машины

Понятие обучения, так же как понятие мышления, сознания и другие психологические понятия, трудно точно определить приемлемым способом. Приблизительная формулировка может быть выражена следующим образом. Предположим, что какой-то организм или машина помещаются в какую-то среду или связаны с группой внешних сигналов и что существует мера «успеха» или приспособления к обстановке. Кроме того, предположим, что эта мера имеет локальный характер во времени, т. е. что можно измерить успех для периодов времени, коротких по сравнению с продолжительностью жизни организма. Если эта локальная мера «успеха» имеет тенденцию улучшаться со временем для рассматриваемой

группы внешних сигналов, можно утверждать, что организм или машина научились приспосабливаться к этой обстановке в соответствии с выбранной мерой успеха. Обучение приобретает количественное значение в выражении обширности и сложности класса внешних сигналов, к которому машина может приспособиться. Машина, играющая в шахматы с возрастающим числом (частотой) выигрышей в течение срока ее жизни, может быть названа в соответствии с этим определением машиной, способной обучаться играть в шахматы; класс внешних сигналов в данном случае образуют игроки-противники, играющие с машиной, а меру приспособления (успеха) — количество выигранных партий.

Был предпринят ряд попыток сконструировать простые обучающиеся машины. Автор данной статьи сконструировал машину для нахождения пути в произвольном лабиринте, состоящем из 25 квадратов, располагаемых в пяти рядах по пять квадратов в каждом. Перегородки между соседними квадратами помещаются также произвольно в зависимости от желания лица, составляющего лабиринт. «Мышь»¹⁾, представляющая собой постоянный магнетик, помещенная в лабиринт, движется сначала ощупью, пробуя пройти и делая ошибки; она наталкивается на различные перегородки и идет вслепую до тех пор, пока не находит дорогу к «кормушке». Когда мышь пускается в другой раз, она направляется без ошибок и неправильных ходов из любой части лабиринта, которую она прошла в первый раз прямо к «кормушке». Помещенная в другую часть лабиринта, она идет вслепую до тех пор, пока не попадает в уже освоенную часть пути, и отсюда идет уже прямо к цели. При этом она добавляет информацию об этой части лабиринта к уже имеющейся информации в своей памяти, и если она снова будет помещена в ту же точку, она прямо направится к цели. Таким образом, мышь, помещаемая в различные неосвоенные части лабиринта, постепенно накапливает полную информацию и может безошибочно достигнуть цели из любой точки лабиринта.

Если затем изменить лабиринт, мышь сначала будет делать попытки пробегать по старому пути, но, наткнувшись на перегородку, постарается найти другое направление, пересматривая свою память до тех пор, пока она не достигнет цели по какому-нибудь другому пути. Таким образом, когда условия задачи меняются, она может забыть старое решение.

В действительности мышь приводится в движение электромагнитом, движущимся под лабиринтом. Движением электромагнита управляет релейное устройство, содержащее около 110 реле, обра-

¹⁾ В первом варианте этой машины (см. наст. сборник, стр. 223—231) вместо «мыши» в лабиринте перемещался специальный шуп, приводимый в движение двумя релейными следящими системами. — *Прим. перев.*

зующих память и вычислительную схему, что-то вроде цифровой вычислительной машины.

Прибор, решающий лабиринтную задачу и являющийся весьма примитивным, способен: 1) решать задачи по методу испытаний и ошибок, 2) повторять решения без ошибок, 3) добавлять новую информацию и устанавливать ее соотношение с частным решением, 4) забывать решение, если оно стало неприменимым.

Другим подходом к созданию обучаемой машины явилось использование цифровой вычислительной машины с соответствующей программой.

А. Е. Оттингер разработал две программы обучения для вычислительной машины типа «ЭДЗАК» в Кембриджском университете (Англия). В первой из них машина выполняла две роли: роль обучаемой машины и окружающей среды. Окружающая среда представляет собой абстракцию ряда магазинов, в которых могут быть куплены различные товары; в различных магазинах хранятся различные виды товаров. Обучаемая машина сталкивается с необходимостью запомнить, где могут быть куплены те или иные товары. Начиная работать без предварительных знаний, где можно отыскать данный товар, она беспорядочно ищет среди магазинов до тех пор, пока не находит намеченное. Когда ей это удастся, она отмечает в памяти, где найден товар, и следующий раз прямо идет в тот «магазин», где она уже до этого «покупала». Дополнительной особенностью этой программы было введение в обучаемую машину немного «любопытства». Когда ей удастся найти j -й товар в каком-то определенном «магазине», она интересуется, содержится ли в данном «магазине» $(j+1)$ -й и $(j-1)$ -й товары и отмечает этот факт в своей памяти.

Вторая программа обучения, описанная Оттингером¹⁾, моделирует нечто вроде поведения животных при выработке условных рефлексов. В машину может быть введен стимул различной интенсивности, представляемый в виде целых чисел, подаваемых на вход. На этот стимул машина может реагировать несколькими различными способами, которые выражаются соответствующими целыми числами, выдаваемыми на выходе. Наблюдая ответ, оператор может выражать машинесвое одобрение или неодобрение, вводя в машину в соответствующий момент некоторое третье целое число. В начале работы реакции на стимулы выбираются случайно. Выражение одобрения увеличивает вероятность правильности предыдущего ответа, указание неодобрения уменьшает эту вероятность. После того как машина «научится» выдавать определенный ответ, подержанный одобрением, величина стимула, необходимого для

¹⁾ См. Обучение цифровой вычислительной машины, *Успехи матем. наук*, II, вып. 5 (7), (1956), 153—160. — Прим. перев.

появления этого ответа, уменьшается, т. е. увеличивается вероятность его появления. Наконец, предусмотрено регулярное уменьшение этой вероятности, когда за ответом не следует одобрения.

Дальнейшее усложнение программ этого рода ограничивается только емкостью памяти вычислительной машины, энергией и изобретательностью инженера, составляющего программу. К сожалению, элементарные команды, имеющиеся в большинстве вычислительных машин, плохо приспособлены для логических требований программ обучения, и, следовательно, машины используются неэффективно. Может потребоваться двенадцать или более команд для представления логически простой и часто используемой операции, встречающейся в программе обучения.

Другой тип обучаемой машины был сконструирован Д. В. Хегельбергером¹⁾ для игры с человеком «в монетку».

На пульте управления машины находятся кнопка для пуска, две лампочки-сигнала, обозначенные + и —, и ключ-переключатель, положения которого также обозначены через + и —. Начиная играть с машиной, игрок выбирает + или — и затем нажимает пусковую кнопку. Машина зажигает один из двух сигналов. Если машина зажигает сигнал, соответствующий выбранному игроком, она выигрывает, в противном случае выигрывает игрок. Когда игра заканчивается, игрок регистрирует (вводит в машину) сделанный им выбор соответствующим поворотом ключа-переключателя.

Машина сконструирована таким образом, чтобы анализировать определенные закономерности в последовательности выборов игрока и старается извлечь выгоду из этих закономерностей, когда она их находит. Например, некоторые игроки имеют тенденцию, если они победили в одном туре, сыграть таким же образом и снова выиграть. Машина учитывает такую ситуацию и, когда выявляется такая тенденция, играет так, чтобы выиграть самой. Если такого рода положений не встречается, машина играет, выбирая ходы случайно.

Было установлено, что машина выигрывает 55—60% игр, тогда как, играя по случайному выбору или играя против партнера, который играет только случайно, она может выиграть только 50% игр. Человеку трудно дать беспорядочную последовательность плюсов и минусов (50% выигрышей обеспечиваются ему самой теорией игры) и еще труднее обыграть машину, наводя ее на ложную закономерность и затем меняя ее.

Вторая машина для игры в «монетку» была сконструирована автором данной статьи. Он придерживался той же самой тактики

¹⁾ Хегельбергер, СИИР — автомат, экстраполирующий последовательности, Кибернетический сборник, I, ИЛМ., 1960, стр. 275—289.
— *Прим. перев*

игры, но использовал другой критерий для принятия решения, когда лучше играть случайно или в соответствии с явным планом. После длительных споров о том, какая из двух машин может победить другую, и бесплодных попыток решить математическим путем весьма сложную статистическую задачу, возникающую при совместной работе обеих машин, пришли к выводу, что эту задачу следует решить экспериментальным путем. Была сконструирована третья небольшая машина, на которую были возложены функции посредника и которая должна была передавать информацию относительно готовности машин сделать ход и о сделанном выборе. Эти три машины были соединены вместе и играли несколько часов, причем зрители заключали небольшие пари по ходу игры и подбадривали машины громкими криками. В результате выяснилось, что меньшая из двух машин, действовавшая быстрее, постоянно выигрывала у большей в отношении приблизительно 55 к 45.

Еще один тип обучаемой машины, в большей степени отличающийся от описанных выше, был сконструирован У. Россом Эшби¹⁾, который называл эту машину гомеостатом. Гомеостезис — слово, введенное Вальтером Б. Кэнноном, означает способность животных с помощью обратной связи стабилизировать такие биологические переменные, как температура тела, химический состав крови и т. д. Прибор Эшби является разновидностью самостабилизирующейся сервосистемы. Первая модель гомеостата состояла из четырех взаимосвязанных сервосистем. Соединение между ними осуществлялось при помощи четырех ламельных переключателей и сопротивлений, припаянных к ламелям. Таким образом, изменение баланса сопротивлений в петле одной из сервосистем влияет на другие три петли в зависимости от величин сопротивлений, включенных ламельным переключателем, связанным с этой петлей. Если одна из сервосистем выводится из равновесия, срабатывает соответствующее ограничительное реле, вызывающее перемещение соответствующего ламельного переключателя на один шаг. Как правило, такая связанная сервосистема с четырьмя степенями свободы и случайными значениями взаимных и собственных коэффициентов усиления не будет устойчивой. Если это так, то начинает работать один или несколько ламельных переключателей, изменяющих значения сопротивлений, и их новая комбинация должна дать новую группу значений коэффициентов усиления. Если эта совокупность значений вновь оказывается нестабильной (т. е. образует неустойчивую сервосистему), происходит новое переключение ламельных переключателей. Процесс будет повторяться до тех пор, пока не будет найдено устойчивое состояние. Вели-

¹⁾ См. Эшби У. Р., Введение в кибернетику, ИЛ, М., 1959.—
Прим. перев

чины сопротивлений, связанных с ламельными переключателями, выбирались случайно (путем использования таблицы случайных чисел). Были предусмотрены средства для введения произвольных изменений или ограничений в работу отдельных сервосистем гомеостата. Например, их коэффициенты связи могли быть заменены на обратные; два из них могли быть равными, один мог иметь фиксированную величину и т. д. При всех этих условиях гомеостат был способен найти устойчивое состояние, при котором все сервосистемы находились в покое. Если считать, что задачей устройств является стабилизация сервосистем и что влияние окружающей среды реализуется путем введения различных изменений и ограничений, осуществляемых оператором, можно утверждать, что гомеостат приспособился к окружающей его среде.

Некоторые особенности гомеостата весьма интересны в качестве основы для создания обучаемых машин и для моделирования работы мозга. Этот прибор, как представляется, делал гораздо больше, чем замышлялось его изобретателем. Например, он мог стабилизироваться в условиях, не предусмотренных при конструировании. Особенно интересно использование случайно выбираемых сопротивлений; это в какой-то степени напоминает случайные связи между нейронами в мозгу. Эшби считает, что общий принцип, воплощенный в гомеостате, который он называет ультраустойчивостью, может лежать в основе анализа работы нервной системы животных. По мнению Эшби, одной из основных трудностей, мешающих непосредственному применению этой теории, является тот факт, что время, необходимое для нахождения стабильного решения, увеличивается приблизительно экспоненциально с ростом числа степеней свободы. При числе степеней свободы всего лишь около двадцати для стабилизации системы потребовалось бы время, в несколько раз превосходящее длительность человеческой жизни. Попытка преодолеть эту трудность ведет к необходимости создания такой сложной конструкции, что чрезвычайно трудно решить, насколько эффективно она будет работать. Наш математический аппарат недостаточно совершенен для решения этой задачи, и поэтому в высшей степени желательна дополнительная экспериментальная работа в этом направлении.

Самовоспроизводящиеся машины

В «Эреване» процесс воспроизведения в обществе рисуется как вид симбиотического сотрудничества между человеком и машиной; машины используют людей как посредников для производства новых машин, когда старые изнашиваются. Роль человека напоминает здесь роль пчелы в опылении цветов. Недавно фон Нейман произвел теоретическое исследование задачи настоящего

самовоспроизведения машин и создал две различные математические модели таких машин.

Первая из этих моделей может быть охарактеризована следующим образом. Машины в модели сконструированы из небольшого числа (порядка двадцати) типов элементарных блоков. Эти блоки выполняют сравнительно простые функции, например функции стоек (для структурных целей); элементарных логических элементов, аналогичных упрощенным реле или нейронам (для вычисления); чувствительных элементов (для обнаружения наличия других элементов); соединяющих устройств, аналогичных паяльнику (для соединения элементов между собой) и т. д. Из этих блоков могут быть «сделаны» различные типы машин. В частности, можно сконструировать некоторую универсальную машину, аналогичную универсальной вычислительной машине Тьюринга. В универсальную конструирующую машину может быть подана последовательность команд, аналогичных программе для цифровой вычислительной машины, которые в виде соответствующего кода программируют конструирование любой другой машины, которая может быть создана из элементарных блоков. Универсальная конструирующая машина затем производит подбор необходимых элементов в своем окружении и конструирует машину в соответствии с информацией, содержащейся на ленте. Если команды, подаваемые в универсальную конструирующую машину, являются описанием самой универсальной конструирующей машины, то машина воспроизводит себе подобную машину и является самовоспроизводящей машиной, за исключением того, что получившаяся копия не содержит программы. Если добавить к универсальной машине лентокопирующее устройство и сравнительно простое управляющее устройство, получаем настоящую самовоспроизводящуюся машину. Теперь команды описывают исходную универсальную машину с добавлением устройства, копирующего ленту, и управляющего устройства. Первая операция этой машины будет заключаться в воспроизведении собственно машины. Управляющее устройство направляет затем ленту с командами в копирующее устройство и помещает полученную копию во вторую машину. Наконец, оно включает вторую машину, которая начинает читать команды с этой ленты и изготавливать третью копию, и так до бесконечности.

Совсем недавно фон Нейман перешел от этой механической модели к более абстрактной самовоспроизводящейся структуре, в основу которой были положены двумерные матрицы элементарных ячеек. Каждая ячейка имеет сравнительно простое внутреннее строение, фактически около 30 возможных внутренних состояний, и непосредственно сообщается лишь с четырьмя соседними с ней ячейками. Состояние ячейки в следующий такт зависит только от текущего состояния ячейки и состояний четырех соседних с нею

ячеек. Путем соответствующего выбора этих переходов состояний можно создать структуру, реализующую в некотором роде самовоспроизведение. Группа смежных ячеек может действовать как единый организм и воздействовать на соединение ячейки, преобразуя эту группу клеток в аналогичное себе устройство.

Эта вторая модель охватывает многие до некоторой степени посторонние проблемы, как-то: определение местонахождения, узнавание и размещение блоков, которые были рассмотрены в первой модели, и, следовательно, приводит к более простой математической формулировке. Кроме того, она имеет сходство с некоторыми химическими и биологическими проблемами, как-то: рост кристаллов и воспроизведение генов, тогда как первая модель ближе к проблемам самовоспроизведения животных.

В процессе рассмотрения этих двух моделей возникает интересное понятие критической сложности, необходимой для самовоспроизведения. При этом ясно, что воспроизводить себя могут лишь достаточно сложные машины.

Фон Нейман считает, что машина, обладающая этим свойством, должна иметь десятки тысяч элементов. Менее сложные машины могут конструировать лишь машины, более простые, чем они сами, тогда как более сложные машины способны обладать своего рода «эволюционными» улучшениями, ведущими к созданию «организмов», более сложных, чем их создатели.

Обращение к читателю

Мы надеемся, что вышеприведенные примеры неарифметического использования вычислительных машин будут стимулировать исследовательскую работу читателей в этой области. Проблема изучения работы мозга и проектирования машин для моделирования его работы является, безусловно, одной из самых важных и сложных научных проблем в настоящее время. Требуется выяснения бесчисленное количество вопросов, охватывающих область как экспериментальной и исследовательской работы, с одной стороны, так и чисто математических исследований, с другой. Можно ли сконструировать машину с локально случайными связями? Возможна ли организация машины по иерархии уровней, как представляется сейчас устройство мозга, с постепенным прогрессирующим ее обучением? Можно ли составить программу для цифровой вычислительной машины таким образом, чтобы (в конечном счете) 99% команд, которым она повинует (а не несколько процентов, как в существующих в настоящее время программах), писались бы самим вычислительным устройством? Может ли быть спроектирована саморемонтирующаяся машина, которая находит неисправности и устраняет их при помощи собственных элементов (включая элемен-

ты в отделении для запасных частей)? Что добавляет элемент случайности к машине Тьюринга? Могут ли быть созданы и управляться при помощи вычислительных устройств манипулирующие и чувствительные приборы, функционирующие аналогично рукам и глазам человека? Может ли быть фактически изготовлена какая-нибудь из самовоспроизводящихся моделей фон Неймана? Может ли быть сформулирована более удовлетворительная теория обучения? Может ли быть сконструирована машина, которая будет проектировать другие машины, если будут заданы лишь их общие функциональные характеристики? Какой должна быть по-настоящему хорошая система команд в универсальных цифровых вычислительных машинах для решения неарифметических задач? Как должна быть организована память вычислительной машины, чтобы обеспечить обучение и вспоминание по ассоциации, аналогично тому, как это делает человеческий мозг?

Предлагаем читателю задуматься над этими вопросами из области теории автоматов. Это — неподнятая целина для ученых. Речь идет не о разработке старых месторождений, а об открытии новых богатых жил и, пожалуй, в некоторых случаях просто о том, чтобы подобрать самородки, лежащие на поверхности.

ЛИТЕРАТУРА

- Ashby W. R., Design for a brain, New York, Wiley, 1951; русский перевод: Эшби У. Р., Конструкция мозга, ИЛ, М., 1962.
- Berkeley E. C., Giant brains, or machines that think, New York, 1949.
- Butler S., Erewhon and erewhon revisited, New York, 1927.
- Diebold J., Automation, New York, Van Nostrand, 1952.
- Householder A. S., Landahl H. D., Mathematical biophysics of the central nervous system, Bloomington, 1945, p. 103—110.
- Kleene S. C., Representation of events in nerve nets and finite automata, Rand Corporation Memorandum RM-704, 1951; русский перевод в сб. Автоматы, ИЛ, М., 1956.
- McCallum D. M., Smith J. B., Mechanized reasoning, *Electronic Engineering*, April, 1951.
- McCulloch, W. S., Pitts W., A logical calculus of the ideas immanent in nervous activity, *Bull. Math. Biophysics*, 5 (1943), 115—133; русский перевод в сб. Автоматы, ИЛ, М., 1956.
- McCulloch W. S., The brain as a computing machine, *Electrical Engineering*, June (1949).
- Meszar J., Switching systems as mechanical brains, *Bell Labs. Record*, 31 (1953), 63—69.
- Oettinger A., Programming a digital computer to learn, *Phil. Mag.*, 43, December (1952), 1243—1263; русский перевод: *Успехи матем. наук*, 11 (1956), вып. 5 (71).

- P e a s e W., An automatic machine toll, *Scientific American*, 187, September (1952), 101—115.
- S h a n n o n C. E., Presentation of a maze-solving machine, Transactions of the Eighth Cybernetics Conference, Josiah Macy, Jr. Foundation, New York, 1952, p. 173—180 (см. наст. сборник, стр. 223).
- S h a n n o n C. E., Programming a computer for playing chess, *Phil. Mag.*, 41, March (1950), p. 256—275 (см. наст. сборник, стр. 192).
- S t r a c h e y C. S., Logical or non-mathematical programmes, *Proc. of the Assn. for Computing Machinery*, Toronto (1952), p. 46—49.
- T u r i n g A. M., Computing machinery and intelligence, *Mind*, 59 (1950), 433—460; русский перевод: Тьюринг А., Может ли машина мыслить, Физматгиз, М., 1960.
- T u r i n g A. M., On computable numbers with an application to the Entscheidungsproblem, *Proc. Lond. Math. Soc.*, 24 (1936), 230—265.
- v o n N e u m a n J., The general and logical theory of automata from cerebral mechanisms in behavior, New York, Wiley, 1951, p. 1—41; русский перевод: Дж. фон Нейман, Общая и логическая теория автоматов. Приложение в вышеупомянутой книге Тьюринга.
- v o n N e u m a n J., Probabilistic logics, California Institute of Technology, 1952; русский перевод в сб. «Автоматы», ИЛ, М. (1956).
- W i e n e r N., Cybernetics, New York, Wiley, 1948; русский перевод: Винер Н., Кибернетика, Советское Радио, М., 1958.

Столетиями философы и ученые спорили, является ли мозг человека функционально машиной. Можно ли сконструировать машину, способную «мыслить»? За последние десятилетия было построено несколько электронных универсальных вычислительных машин, которые оказались способными осуществлять процесс, в большой степени напоминающий процесс мышления. Эти новые вычислительные машины были построены сначала исключительно для математических расчетов. Они выполняют автоматический длинную последовательность сложений, умножений и других арифметических операций со скоростью тысяч операций в секунду. Основной принцип устройства этих машин придает им такую универсальность и гибкость, что они могут быть приспособлены для работы с элементами символической логики, представляющими слова, высказывания и другие понятия.

Одной такой возможностью, которая уже исследовалась с различных точек зрения, является перевод с одного языка на другой при помощи машины. Ближайшей целью ставится не окончательная литературная редакция текста, а только пословный перевод, определяющий значение слов. Вычислительные машины могут также служить для многих других задач полумеханического, полумыслительного характера, таких, как: конструирование электрических фильтров и релейных схем, составление расписаний рейсов в загруженных аэропортах, коммутирование телефонных разговоров наиболее эффективным образом через ограниченное число каналов.

Некоторые возможности в этом направлении могут быть проиллюстрированы путем приспособления вычислительной машины для достаточно хорошей игры в шахматы. Эта проблема, конечно, не имеет серьезного практического значения сама по себе, но она была поставлена с серьезными целями. Исследование задачи игры в шах-

¹⁾ Shannon C., Chess playing machine *The world of mathematics*, 4 (1956), 2124. [Первоначально статья была опубликована в журнале *Sci. Amer.*, 182, 2, 1950, 48. — Прим. ред.]

маты на машине имеет в виду создание технических устройств, которые могут быть использованы для многих практических применений.

Игра в шахматы является идеальной для начальных попыток в этом направлении по нескольким причинам. Она четко определена как допустимыми операциями (шахматные ходы), так и конечной целью (мат). Она не является ни настолько простой, чтобы быть



Рис. 1. Шахматный автомат XVIII века действительно управлялся человеком, спрятанным внутри.

тривиальной, ни слишком трудной для получения удовлетворительного решения. Машина, играя в шахматы против человека, позволяет выяснить возможности ее применения в решении такого рода задач.

Уже имеется значительная литература по описанию машин, играющих в шахматы. В конце XVIII и начале XIX века венгерский изобретатель Вольфганг фон Кемпелен поразил Европу устройством, которое известно под названием Маельзельского (Maelzel's) шахматного автомата. Автор демонстрировал автомат, путешествуя по континенту. Вскоре появилось большое число статей, объяснявших его действие, включая аналитическую статью Эдгара Аллана По. Большинство авторов заключало вполне правильно, что автомат управляется шахматным мастером, скрытым внутри. Несколько лет спустя был вскрыт истинный принцип его работы (см. рис. 1).

Более честная попытка создать машину, играющую в шахматы, была сделана в 1914 г. испанским изобретателем Л. Торресом-и-Кведо (L. Torres y Quevedo), который построил автомат, разыгрывающий эндшпиль — король и ладья против короля. Машина,

играющая за партнера с королем и ладьей, должна была форсировать мат в несколько ходов, как бы ни играл противник — человек. Так как можно дать точную последовательность правил для выбора удовлетворительных ходов в таком окончании, проблема сравнительно проста, но сама идея такой машины была для своего времени прогрессивной.

Электронную вычислительную машину можно заставить довольно сильно играть в шахматы. Для того чтобы показать это, лучше всего начать с общего описания вычислительной машины и ее действия.

Электронная универсальная вычислительная машина является чрезвычайно сложным устройством, содержащим несколько тысяч электронных ламп, реле и других элементов. Однако основные принципы, положенные в основу ее действия, совершенно просты.

Машина имеет четыре главные части:

- 1) арифметическое устройство,
 - 2) управляющее устройство,
 - 3) устройство для запоминания и хранения чисел (числовая память),
 - 4) устройство для запоминания программ (программная память).
- (В некоторых конструкциях устройства для запоминания чисел и программ устроены одинаково.)¹⁾

Действие машины при вычислениях в точности аналогично работе человека, проводящего серию вычислений на обычной настольной вычислительной машине. Арифметическое устройство соответствует настольной машине, управляющее устройство — человеку, выполняющему вычисления, числовая память — рабочему листу, на котором записываются промежуточные и конечные результаты, а в программной памяти хранится алгоритм, описывающий ту последовательность операций, которая должна быть выполнена.

В электронных вычислительных машинах числовая память состоит из большого числа «ячеек», каждая из которых способна хранить число. Для того чтобы решить задачу на машине, необходимо записать в ячейки числа, соответствующие входящим в расчетные формулы величинам, и затем написать программу, которая указывает машине, какие арифметические операции должны быть выполнены и куда должны быть записаны результаты. Программа состоит из последовательности «команд», каждая из которых описывает элементарную операцию. Например, типичной командой служит выражение:

Сл. 372, 451, 133.

¹⁾ В большинстве современных машин в любой ячейке памяти может находиться либо число, либо команда.— *Прим. перев.*

Оно означает: сложить число, хранящееся в ячейке 372, с числом из ячейки 451 и получившуюся сумму записать в ячейку 133. Другой тип команд заставляет машину сделать выбор продолжения решения. Например, команда:

Ср. 291, 118, 345

означает: сравнить содержимое ячеек 291 и 118 и, если число в ячейке 291 больше числа в ячейке 118, перейти к выполнению следующей по порядку команды; если нет, перейти к выполнению команды, хранящейся в ячейке 345. Этот тип команд позволяет машине сделать выбор из двух возможностей в зависимости от результатов предыдущих вычислений. «Словарь» электронной машины может включать около 30 различных типов команд.

После того как в машину введена программа и исходные числа, требуемые для вычислений, расположены в числовой памяти, она автоматически производит вычисления. Конечно, такая машина наиболее удобна для решения задач, включающих колоссальное число отдельных вычислений, которые было бы чрезвычайно затруднительно провести вручную.

Задача постановки игры в шахматы на вычислительной машине может быть разделена на три основных этапа.

1. Должен быть выбран определенный код для представления числами шахматной позиции и шахматных фигур.

2. Должна быть выбрана некоторая стратегия для выбора хода.

3. Эта стратегия должна быть переведена в последовательность элементарных команд машины или программу.

На рис. 2 показан код для шахматной доски и фигур. Каждый квадрат имеет номер, состоящий из двух чисел: первое число соответствует строке или горизонтали, а второе — столбцу или вертикали. Разным шахматным фигурам приписаны различные числа: пешке 1, коню 2, слону 3, ладье 4 и т. д. Белым фигурам соответствуют положительные числа, а черным отрицательные. Положение всех фигур на доске может быть представлено последовательностью 64 чисел, причем пустому квадрату соответствует 0. Таким образом, любая шахматная позиция может быть записана как ряд чисел и введена в числовую память машины.

Любой ход определяется заданием номера поля, на котором стояла фигура, и номером поля, на который она пошла. Обычно этих двух номеров бывает достаточно для описания хода, но нужно учитывать специальный случай превращения пешки в более ценную фигуру, когда необходимо третье число. Во всех других случаях третье число равно нулю. Следовательно, ход коня с поля 01 на 22 кодируется как 01, 22, 0. Продвижение пешки с 62 на 72 и превращение ее в ферзя записывается как 62, 72, 5.

Следующей основной задачей является выбор игровой стратегии. Необходимо найти алгоритм для выбора достаточно хорошего хода в любой заданной позиции. Это — самая трудная часть задачи. Составитель программы может использовать принципы хорошей игры, которые были разработаны лучшими шахматными мастерами. Эти эмпирические принципы вносят некоторый порядок

7	-4	-2	-3	-5	-6	-3	-2	-4
6	-1	-1	-1	-1	-1	-1	-1	-1
5								
4								
3								
2								
1	1	1	1	1	1	1	1	1
0	4	2	3	5	6	3	2	4
	0	1	2	3	4	5	6	7

Р и с. 2. Кодирование начальной позиции для шахматной машины изображено в виде шахматной доски. Каждый квадрат может быть определен двумя числами: одно задает горизонталь, а второе — вертикаль. Фигурам также соответствуют числа.

в нагромождение возможных вариантов игры. Даже огромные скорости электронных вычислительных машин совершенно недостаточны для ведения абсолютно точной игры в шахматы путем просчета всех возможных вариантов игры до конца. В типичной шахматной позиции имеется около 32 возможных ходов с 32 допустимыми ответами, это создает 1024 варианта на один ход. В большинстве шахматных партий делается 40 или больше ходов с каждой стороны. Следовательно, общее число вариантов в средней игре составляет около 10^{120} . Машина, вычисляющая один вариант в одну миллионную долю секунды, будет работать 10^{95} лет, прежде чем выберет первый ход!

Другие методы для ведения абсолютно точной шахматной игры представляются такими же нереальными. Следовательно, стремясь к повышению качества игры в шахматы на машине, приходится примириться с тем обстоятельством, что сделанный ею ход не

обязательно будет наилучшим. Таким же образом, конечно, действует и любой шахматист, ибо никто не умеет играть в шахматы абсолютно точно.

При выработке некоторой стратегии для машинной игры в шахматы необходимо установить метод численной оценки любой данной позиции. Шахматист, глядя на доску, может оценить, какая сторона имеет преимущество. Однако его оценка будет грубо количественной. Он может сказать: «Белые имеют ладью за слона, преимущество примерно в две пешки», или: «Черные имеют достаточную мобильность за пожертвованную пешку». Эти суждения базируются на долгом опыте и обобщениях принципов, почерпнутых из шахматной литературы. Например, известно, что ферзь стоит примерно девять пешек, ладья около пяти, а слон или конь около трех. В качестве первого грубого приближения для оценки позиции может служить просто сумма сил с каждой стороны, измеренная в терминах пешечных единиц. Имеются, однако, другие черты шахматной позиции, которые должны учитываться: мобильность и расположение фигур, слабая защита короля, пешечная структура и т. д. Таким факторам шахматной позиции тоже необходимо приписать определенные веса и включить в оценочную функцию, именно здесь необходимо использовать знание и опыт шахматных мастеров.

Если подходящий метод численной оценки позиции найден, как надо выбирать ход? Простейший процесс заключается в том, что просматриваются все возможные ходы в данной позиции и выбирается один, который дает наивысшую оценку. Так как шахматист обычно просматривает далее, чем на один ход вперед, можно учитывать возможные ответы противника на каждый ожидаемый ход. Предполагая, что ответ противника должен быть одним из ходов, дающих наилучшую оценку позиции с его точки зрения, следует выбрать ход, который дает нам больше всего после наилучшего ответа партнера. К сожалению, быстроедействие современных электронных машин допускает просмотр вперед только на два хода с каждой стороны, так что при использовании такой стратегии на машине она будет играть неважно сравнительно с человеком. Хорошие шахматные мастера часто проводят комбинации на четыре, пять ходов, а чемпионы мира просматривают некоторые варианты на 20 ходов вперед. Это возможно только потому, что варианты, которые просматривались, строго отбирались. Они не рассматривали всех направлений развития игры, а только самые важные.

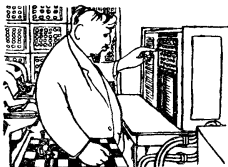
Глубина анализа и выбора, достигнутая шахматными мастерами при проверке различных вариантов, была изучена экспериментально датским шахматистом и психологом А. Д. де Гроотом (A. D. De Groot). Он показывал различные типовые позиции шахматным мастерам и просил их выбрать наилучший ход, описывая

вслух ход анализа позиции, который они обычно проводят при игре. Таким образом можно было определить число и глубину исследуемых вариантов. В одном типичном случае шахматист проверил до 16 вариантов, разнящихся по глубине от одного хода за черных до пяти ходов за черных и четырех за белых. Всего им было рассмотрено 44 позиции.

Ясно, что было бы чрезвычайно желательно улучшить стратегию для машины включением такого процесса анализа и выбора. Конечно, нельзя заходить слишком далеко в этом отношении. Исследование одного направления развития игры на 40 ходов вперед настолько же плохо, как и исследование вариантов только на два хода. Подходящим компромиссом было бы исследование только важнейших возможных вариантов, таких, как форсированные варианты взятия фигур и основные угрозы, и продолжение их исследования настолько далеко, чтобы проверить каждый такой вариант до полной ясности. Вполне возможно установить некоторые грубые критерии для выбора важнейших вариантов, конечно, не так эффективно, как это делает шахматист, но достаточно для того, чтобы ощутимо уменьшить число вариантов и, следовательно, позволить рассматривать достаточно глубоко выбранные варианты.

Последняя задача состоит в сведении стратегии к последовательности команд, переводящих ее на язык машины. Это сравнительно техническая, но тяжелая работа, и здесь будут указаны только некоторые общие ее черты. Вся программа состоит из девяти подпрограмм и основной программы, которая обращается к подпрограммам в случае необходимости. Шесть из этих подпрограмм имеют дело с движением различного рода фигур. Они указывают машине дозволенные ходы этих фигур. Еще одна подпрограмма позволяет машине делать ход в «уме» без действительного передвижения фигур, т. е. из данной позиции, сохраняемой в памяти машины, она создает позицию, получающуюся в результате некоторого хода. Седьмая подпрограмма дает возможность вычислительной машине составить список ходов, возможных в данной позиции. Последняя подпрограмма вычисляет оценочную функцию для любой заданной позиции. Основная программа координирует работу остальных подпрограмм. Сначала подпрограмма 7 составляет список возможных ходов, пользуясь в свою очередь подпрограммами 1—6 для того, чтобы определить, куда может пойти данная фигура. Затем основная программа оценивает результирующие позиции при помощи подпрограммы 8 и сравнивает результаты с помощью процесса, описанного выше. После сравнения всех исследуемых вариантов выбирается тот ход, который дает наилучшую оценку. Этот ход затем печатается машиной в нормальной шахматной нотации.

Можно надеяться, что по описанной выше программе машина будет играть довольно сильно и по скорости сравнимо с человеком. Машина имеет ряд очевидных преимуществ по сравнению с человеком-шахматистом.



Р и с. 3. Иллюстрация неизбежного преимущества человека перед машиной. В верхней части человек проигрывает машине. В середине, разозлившись, он меняет команды в программе. Внизу человек выигрывает.

1) Она делает конкретные вычисления с гораздо большей скоростью.

2) Ее игра свободна от ошибок, не считая программных слабостей, в то время как человек часто делает простые и очевидные просмотры.

3) Машина никогда не будет лениться, никогда не поддастся соблазну сделать ход инстинктивно без надлежащего анализа позиции.

4) У машины нет «нервов», следовательно, она не будет переоценивать свою позицию или наоборот — недооценивать свои шансы.

В противовес этим качествам ум человека обладает гибкостью, воображением и способностью к обучению.

В некоторых случаях машина может победить человека, задавшего ей стратегию игры. Но в одном смысле он все же остается сильнее, — зная стратегию игры машины, человек может применить ту же самую тактику, но продолжив ее на большую глубину. При этом, возможно, ему потребуется несколько недель для выбора хода, а машина вычислит его в несколько минут. При одинаковом количестве времени на «обдумывание» быстрое действие, выдержка и высокая

точность машины позволят ей успешно играть с человеком. Однако, сильно разозлившись, создатель программы может легко ослабить игровое мастерство машины, изменив программу таким образом, чтобы уменьшилась глубина исследования (см. рис. 3).

Эта идея была выражена в карикатурах на страницах журнала (The Saturday Evening Post) год назад.

Если поступать так, как только что описывалось, машина будет делать один и тот же ход при каждом повторении позиции. Если противник будет повторять свои ходы, то полностью повторится вся партия. Однажды выиграв партию, он может выигрывать каждый раз, применяя одну и ту же тактику и пользуясь тем обстоятельством, что машина в некоторой позиции сделала слабый ход. Одним из способов, при помощи которых можно «разнообразить» игру машины, является введение в нее случайных элементов. Например, если машине представляется возможным сделать два или более ходов в данной позиции, оценки которых примерно одинаково хороши «согласно ее вычислениям», она выбирает один из них случайным образом. Поэтому, придя вторично к той же самой позиции, машина может сделать другой ход.

Случайные элементы могут быть введены также в начале игры. Желательно иметь много стандартных дебютов с различными вариантами (до нескольких сотен), которые хранятся в памяти машины. Начиная с первого хода, до тех пор пока противник не отклонится от стандартных ответов или машина не исчерпает последовательности ходов, записанных в памяти, — она будет играть по теории. Это нельзя назвать обманом, так как таким же образом шахматные мастера разыгрывают дебют.

Заметим, что в определенных пределах машина будет играть блестяще; она с готовностью пойдет на жертвы важнейших фигур, чтобы позже получить преимущество или дать мат, если окончание комбинации находится в пределах ее вычислительных возможностей. Например, в позиции, изображенной на рис. 4, машина быстро найдет мат в три хода с жертвами:

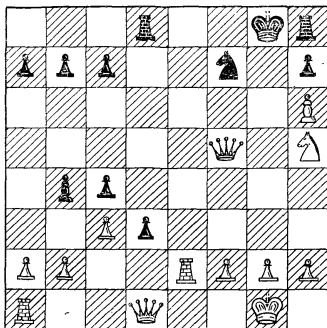
Белые	Черные
1. Лe8+	Л: e8
2. Фg4+	Ф: g4
3. Кf6×	

Комбинации такого типа часто не замечаются в любительских играх.

Основная слабость машины состоит в том, что она не учится на своих ошибках. Единственным путем для улучшения ее игры является усовершенствование программы. Некоторые мысли о создании программы, которая по мере накопления опыта вносила бы улучшения в исходную стратегию, были уже высказаны. И хотя это теоретически возможно, но методы, которые до сих пор предлагались, вряд ли применимы на практике. Одной из возможностей является создание программы, которая будет менять члены и коэф-

фициенты оценочной функции на основании результатов игр, в которых машина участвовала раньше. Для этого могут быть введены небольшие изменения в эти члены и должны быть выбраны их значения, дающие наибольший процент побед.

Главный вопрос, который легче задать, чем на него ответить, это: воспроизводит ли машина, играющая в шахматы, процесс «мышления»? Ответ будет зависеть от того, как определить мышление. Так как сейчас нет общего решения о точном смысле этого термина, на этот вопрос нельзя дать определенный ответ. С точки зрения поведения машина действует так, как будто она мыслит.



Р и с. 4. Задача, которую машина может решить блестяще, дается этой шахматной позицией. Машина пожертвует ладью и ферзя, ценнейшие фигуры, чтобы выиграть следующим ходом.

Всегда считалось, что искусная игра в шахматы требует мыслительных способностей. И если рассматривать мышление с точки зрения внешних проявлений, а не внутреннего метода, машина, несомненно, мыслит.

Мыслительный процесс, по мнению некоторых психологов, существенно характеризуется следующими этапами: различные возможные решения задачи испытываются мысленно или символически без реального воспроизведения; лучшее решение выбирается мысленной оценкой результатов этих экспериментов; затем реализуется решение, найденное таким путем. Нужно заме-

тить, что это почти точное описание того, как действует, играя в шахматы, машина, если заменить слово «мысленно» словами «внутри машины».

С другой стороны, машина делает только то, что ей приказано. Она работает, делая попытки и ошибаясь, но она делает только то, что ей приказал делать автор программы; она делает ошибки потому, что по оценочной функции она находит некоторые варианты нестоящими. Машина находит решение, но ее решения были предусмотрены при выработке стратегии. Короче, машина не выходит за рамки того, что было в нее заложено. Это положение было прекрасно понято Торресом-и-Квеvedо, который по поводу своей машины, разыгрывающей определенное окончание шахматной игры, заметил: «Пределы, внутри которых мышление действительно необходимо, требуют лучшего определения... Автомат может делать многое из того, что в жизни называют мышлением».

СОСТАВЛЕНИЕ ПРОГРАММ ДЛЯ ИГРЫ В ШАХМАТЫ НА ВЫЧИСЛИТЕЛЬНОЙ МАШИНЕ¹⁾

1. Введение

В статье рассматривается задача составления программы для игры в шахматы на современной универсальной вычислительной машине. Эта задача, не имея, возможно, практической ценности, представляет теоретический интерес. Можно надеяться, что ее удовлетворительное решение явится плацдармом для наступления на другие задачи сходной природы, но имеющие большее практическое значение. Вот некоторые возможности в этом направлении.

1) Конструирование фильтров, эквипотенциальных соединений и т. д.

2) Конструирование переключательных схем.

3) Управление распределением телефонных вызовов в зависимости от складывающихся обстоятельств, а не по фиксированной схеме.

4) Выполнение символических (не арифметических) математических операций.

5) Перевод с одного языка на другой.

6) Принятие решений командиром по упрощенным вариантам оперативной обстановки.

7) Оркестровка мелодий.

8) Выработка логических выводов.

Вероятно, разнообразные устройства различной природы для решения перечисленных задач будут созданы в ближайшем будущем. Техника, развитая для современных вычислительных машин, делает решение этих задач не только теоретически возможным, но в некоторых случаях позволяет серьезно рассматривать их с экономической точки зрения.

Универсальные машины имеют следующие преимущества перед обычными вычислительными машинами.

¹⁾ Shannon C., Programming a computer for playing chess, *Phil. Mag.*, 41, ser. 7, № 314 (1950), 256.

Во-первых, они могут выполнять операции не только над числами, но также над шахматными позициями, схемами, математическими выражениями, словами и т. д.

Во-вторых, сам вычислительный процесс включает некоторые общие принципы, своего рода разумный выбор, возможность изменения процесса решения на основе метода испытаний и ошибок, а не только на основе заранее заданных правил.

Наконец, решая указанные задачи, машины могут давать не только два ответа — правильный или неверный, — но непрерывный ряд ответов от лучшего к худшему. Например, машина, которая конструирует хорошие фильтры, а не обязательно наилучшие из всех возможных, может быть признана удовлетворительной. Машина для игры в шахматы наиболее соответствует первой попытке решения таких неарифметических задач, так как: 1) шахматная игра четко определяется, с одной стороны, допустимыми операциями (ходами), а с другой стороны, конечной целью (мат); 2) эта игра ни настолько проста, чтобы быть тривиальной, ни настолько трудна, чтобы нельзя было найти удовлетворительное решение; 3) считается, что обычно игра в шахматы требует «обдумывания» для достижения успеха; следовательно, решение этой задачи заставит нас либо допустить существование механического мышления, либо приведет к дальнейшему ограничению понятия «мышления»; 4) дискретная структура игры хорошо согласуется с цифровой природой современных вычислительных машин.

В настоящее время имеется значительная литература, описывающая машины для игры в шахматы. В конце XVIII и начале XIX века Маельзельский (Maelzel's) шахматный автомат — устройство, изобретенное фон Кемпеленом (von Kempelen), широко демонстрировался в качестве машины, играющей в шахматы. В это время появилось большое количество статей, включая аналитический очерк Эдгара По (озаглавленный «Маельзельский шахматный игрок»), пытающихся объяснить принцип его работы. Большинство авторов совершенно правильно заключали, что автоматом управляет спрятанный шахматист. Однако аргументы, приводящие к такому заключению, часто были неверными. Эдгар По полагал, например, что можно построить машину, которая будет побеждать постоянно, либо машину, которая будет побеждать случайно, и отсюда выводит, что, поскольку автомат не является непобедимым, то, следовательно, он управляется человеком. Ясно, что такое рассуждение неправильно. Для полного знакомства с историей и принципом работы автомата читатель отсылается к серии статей Харкнеса и Беттела в журнале «Чесс ревью» (*Chess Review*) за 1947 год¹⁾.

¹⁾ Harkness and Battel, *Chess Review*, 1947.

Более честная попытка сконструировать машину, играющую в шахматы, была сделана в 1914 году Торресом-и-Квеvedо, который построил автомат для разыгрывания эндшпиля: король и ладья против короля. Машина играла за сторону, имеющую короля и ладью, и форсировала мат в несколько ходов независимо от игры противника. Так как для выбора правильного хода в таком окончании может быть дан точный ряд правил, то задача относительно проста, но идея такой машины была прогрессивна для своего времени.

Положение, которое будет развито в работе, состоит в том, что современные вычислительные машины могут быть успешно использованы для игры в шахматы. Для этого необходимо составить соответствующую программу. Хотя первое приближение, которое здесь дается, по нашему мнению, в основном полно, требуется сделать еще много как в экспериментальном, так и в теоретическом отношении.

2. Общие положения

Шахматная позиция может быть определена следующими параметрами.

- 1) Описанием положения всех фигур на доске.
- 2) Указанием стороны, делающей ход в данной позиции.
- 3) Указанием того, ходили ли раньше короли и ладьи. Это важно, так как, например, после передвижения ладьи право на рокировку в соответствующую сторону теряется.
- 4) Указанием последнего сделанного хода; оно будет нужно, чтобы определить возможность взятия на проходе, так как эта привилегия после одного хода утрачивается.
- 5) Указанием числа ходов, выполненных после последнего движения пешкой или после последнего взятия. Это важно, так как в случае достижения этим числом 50 присуждается ничья. Для простоты будем пренебрегать правилом присуждения ничьей после трехкратного повторения позиции.

В шахматах нет элементов случайности, кроме первоначального выбора, кому из игроков делать первый ход, в противоположность играм в карты, трик-трак и т. д. Далее, в шахматах каждый из партнеров имеет полную информацию о каждом ходе и о всех предыдущих ходах (в противоположность, например военным ситуациям). Из этих двух фактов вытекает¹⁾, что любая шахматная позиция может быть:

- 1) либо выигрышной для белых (это означает, что белые могут выиграть, как бы ни защищались черные);

¹⁾ Von Neumann J., Morgenstern O., Theory of games and economic behavior, 1944.

2) либо ничейной (белые могут добиться по крайней мере ничьей, как бы ни играли черные, но и черные могут добиться не более, чем ничьей, как бы ни играли белые);

3) выигрышной для черных (черные могут выиграть, как бы ни защищались белые).

Это утверждение является для практических целей аналогом теоремы существования. Неизвестно практического метода для определения того, к которой из этих трех категорий принадлежит любая выбранная позиция. Если такой метод появится, шахматы потеряют интерес как игра. Тогда можно будет определить, является ли начальная позиция выигрышной, проигрышной или ничейной для белых, и исход игры между противниками, знающими этот метод, будет полностью определен выбором того, кому первому ходить. Если исходить из того, что начальная позиция ничейная (что подсказывает опыт игр мастеров¹⁾), каждая игра должна закончиться вничью.

Интересно, что небольшое изменение в шахматных правилах дает игру, для которой можно доказать, как теорему, что белые имеют по крайней мере ничью в начальной позиции. Действительно, предположим, что правила игры те же, что и в настоящие шахматы, но игрок может пропустить свою очередь хода. Тогда можно доказать как теорему, что белые при правильной игре добиваются по крайней мере ничьей.

В самом деле, в начальной позиции либо есть выигрышный ход, либо нет. Если есть — сделаем его. Если нет — передадим очередь хода противнику. Черные тогда будут находиться в том же самом положении, что и белые перед первым ходом вследствие зеркальной симметрии начальной позиции²⁾. Так как белые не имели в этой позиции выигрышного хода, теперь такого хода не имеют черные. Следовательно, черные могут самое большее добиться ничьей. Но тогда белые в любом случае имеют минимум ничью.

В некоторых играх существует простая оценочная функция $f(P)$, которая относится к позиции P и величина которой определяет, к какой категории (выигрыш, проигрыш, ничья) принадлежит позиция P . В игре Ним³⁾, например, она может быть определена записью числа спичек в каждой группе в двоичной системе. Эти числа подписываются друг под другом (как для сложения столбиком). Рассматриваются получившиеся столбцы. Если число единиц

1) Матч на первенство мира по шахматам между Капабланкой и Алехиным закончился победой Алехина $+6, -3, =25$.

2) Тот факт, что уменьшается количество ходов (в обычном смысле), после которого присуждается ничья, согласно правилу 50 ходов, не влияет на рассуждения.

3) Hardy J., Wright E., Theory of numbers, Oxford, 1938.

в каждом столбце четно, то позиция проигрышна для игрока, который должен делать очередной ход, и выигрышна в противном случае.

Если для игры может быть найдена такая оценочная функция $f(P)$, то легко сконструировать машину, способную вести игру наилучшим образом. Она никогда не проиграет и не сведет в ничью выигрышную партию, никогда не проиграет ничейную партию и, если противник ошибется, воспользуется этим. Это можно сделать следующим образом: предположим, что

$f(P)=1$ для выигрышной позиции,

$f(P)=0$ для ничейной позиции,

$f(P)=-1$ для проигрышной позиции.

При своем ходе машина подсчитывает $f(P)$ для различных позиций, получающихся из данной каждым возможным ходом. Она выбирает ход (или один из нескольких ходов), дающий максимум величины f . В случае игры в Ним, где такая функция известна, машина, реализующая оптимальную игру, действительно построена¹⁾.

В шахматах, в принципе, возможно играть абсолютно правильно или построить машину, которая будет это делать, следующим образом: в данной позиции рассматриваются все возможные ходы, затем все ходы за противника и т. д. до конца игры (в каждом варианте). Конец игры, согласно правилам, должен наступить через конечное число ходов²⁾ (вспомните 50-ходовое правило ничьей). Каждый из вариантов заканчивается поражением, ничьей или выигрышем. Рассматривая варианты с конца, можно определить, имеется ли форсированная победа, позиция ничейная или проигрышная. Легко показать, однако, что даже при высоких скоростях, которыми обладают современные вычислительные машины, провести такое исследование практически невозможно.

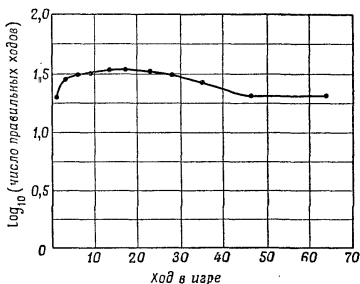
В типовых шахматных позициях бывает около 30 возможных ходов. Это число остается примерно постоянным до тех пор, пока игра не приближается к концу, как показано на рис. 1. Этот график построен на основании данных, полученных де Гроотом³⁾, который определил среднее число допустимых ходов в позиции, обработав большое число партий шахматных мастеров. Таким образом, ход за белых и затем ответ за черных дает около 10^3 вариантов. Обычно шахматные партии длятся около 40 ходов до соглашения

¹⁾ Condon, Tawney and Derr, U. S., Patent 2, 215, 544. Нимотрон-машина, основанная на этом патенте, была построена и экспонировалась Вестингаузом в 1938 г. на Всемирной выставке в Нью-Йорке.

²⁾ Самая длинная игра может продолжаться не более 6350 ходов, если допускать лишь 50 ходов между каждым ходом пешкой или взятием. Длиннейшая напечатанная турнирная партия между мастерами длилась 168 ходов, а самая короткая 4 хода.

³⁾ De Groot A. D., Het Denken van den Shaker, Amsterdam, 1946, 17.

между противниками. Такое определение окончания партии не подходит для наших целей, потому что машина должна вести вычисления до окончания по правилам, а не по соглашению. Однако даже в этом случае подсчет дает около 10^{120} вариантов, которые должны быть проверены в начальной позиции. Машина, работающая со скоростью одного варианта в микро-микро-секунду, потребует более 10^{90} лет, чтобы выбрать первый ход!



Р и с. 1.

Другой (также практически невыполнимый) метод заключается в том, что надо иметь «словарь» всех возможных положений шахматных фигур на доске. Для каждой возможной позиции задается правильный ход (вычисленный указанным выше процессом или подсказанный шахматным мастером). При своем ходе машина просто смотрит на позицию и делает указанный ход. Число возможных позиций достигает порядка $64!/32! (8!)^2 (2!)^6$ или грубо 10^{43} , что, естественно, делает такой метод нереальным.

Ясно, что задача состоит не в том, чтобы построить машину, играющую в шахматы идеально хорошо (что практически совершенно невозможно), и не в том, чтобы просто делать допустимые ходы (что тривиально). Надо научить ее играть искусно, может быть, сравнимо с хорошим шахматистом.

Шахматная стратегия может быть описана как правило для выбора хода в любой данной позиции. Если по этому правилу всегда выбирается один и тот же ход в одной и той же позиции, стратегия в теории игр называется «чистой». Если процесс выбора включает

случайные элементы и не всегда приводит к одному и тому же ходу в одной и той же позиции, то стратегия называется «смешанной».

Приведем простые примеры стратегий.

1. Перенумеруем все возможные ходы в позиции Р согласно некоторому определенному правилу. Будем всегда выбирать ход с номером 1. Это пример «чистой» стратегии.

2. Перенумеруем возможные ходы и будем выбирать номер хода случайным образом. Это пример «смешанной» стратегии.

Обе эти стратегии, конечно, очень бедны и не претендуют на выбор хороших ходов. Наша задача состоит в построении сравнительно хорошей стратегии для выбора очередного хода.

3. Приближенная оценочная функция

Хотя для шахмат простая и точная оценочная функция неизвестна и, вероятно, никогда не будет найдена вследствие сложности природы игры, уже сейчас возможно построить приближенную оценочную функцию. В самом деле, любой хороший шахматный игрок должен уметь делать такую оценку позиции. Оценки базируются на общей структуре позиции, числе и типе белых и черных фигур, пешечной структуре, мобильности и т. д. Эти оценки не совершенны, но чем сильнее игрок, тем лучше его оценки. Большинство правил и принципов правильной игры содержат на самом деле утверждения об оценочной функции.

Приведем некоторые примеры.

1) Сравнительный вес ферзя, ладьи, слона, коня и пешки принимается равным примерно 9, 5, 3, 3, 1 соответственно. Таким образом, при прочих равных условиях (!), если сложить число всех фигур у каждой стороны (с учетом этих коэффициентов), та сторона, у которой эта сумма больше, имеет лучшую позицию.

2) Ладьи следует располагать по открытым линиям. Последнее утверждение — часть более общего принципа о том, что сторона с большей мобильностью при прочих равных условиях имеет лучшую позицию.

3) Отставшие, изолированные или сдвоенные пешки слабы.

4) Открытый король слабее (до перехода игры в окончание).

Эти и аналогичные принципы являются только обобщениями практического опыта и носят статистический характер. Вероятно, любой шахматный принцип может быть опровергнут специальным примером. Однако на основе принципов можно построить грубую оценочную функцию. Примером является следующая функция:

$$f_1(P) = 200(K_p - K_p') + 9(\Phi - \Phi') + 5(L - L') + 3(C - C' + K - K') + \\ + (p - p') - 0,5(D - D' + S - S' + I - I') + 0,1(M - M') + \dots,$$

в которой Кр, Ф, Л, С, К, п обозначают соответственно число белых королей, ферзей, ладей, слонов, коней и пешек на доске. D, S, I есть количество сдвоенных, отставших и изолированных белых пешек соответственно; М равно мобильности белых (измеряемой, скажем, числом возможных ходов белых фигур).

Буквы со штрихом обозначают те же величины для черных.

Коэффициенты 0,5 и 0,1 — просто грубая прикидка автора. Кроме того, имеется еще много разных факторов, которые могут быть учтены¹. Эта формула дается только в целях иллюстрации. Мат искусственно включен в нее приданием королю большей ценности, а именно 200 (что больше, чем максимум того, что могут стоить все остальные слагаемые, вместе взятые).

Можно заметить, что указанная приближенная оценочная функция $f(P)$ имеет более или менее непрерывный спектр возможных значений, в то время как точная оценочная функция имеет только три значения. Так и должно быть. При практической игре позиция может быть «легко выигранной», если, например, игрок имеет преимущество на ферзя, или требует большого труда для выигрыша при пешечном преимуществе. В первом случае имеется много путей для достижения выигрыша, в то время как в последнем требуется точная игра и одна-единственная ошибка часто уравнивает шансы. Теория игр допускает случай безграничной гениальности игрока, однако она же утверждает, что, если не допускать ошибок и использовать малейшее преимущество, можно добиться такого же результата, что и в первом случае. Игра между двумя гениальными игроками мистером А и мистером В проходила бы следующим образом. Они садятся за стол, разыгрывают цвет, а затем мгновение смотрят на доску. После чего либо:

- 1) мистер А говорит: «Я сдаюсь», либо,
- 2) мистер В говорит: «Я сдаюсь», либо,
- 3) мистер А говорит: «Предлагаю ничью», и мистер В отвечает: «Согласен».

4. Стратегия, основанная на оценочной функции

Очень важным обстоятельством, относящимся к оценочным функциям описанного выше типа (и общим принципом шахматной игры), является то, что они приложимы лишь к относительно спокойным позициям. Например, если белые, размениваясь ферзями, забирают черного ферзя, а черные следующим ходом отыгрывают его, бессмысленно вычислять оценочную функцию в тот момент, когда белые временно получили большой выигрыш. Вообще, бесполезно вычислять оценочную функцию общего типа в момент прове-

¹) См. приложение 1.

дения комбинации или серии разменов. Для учета разменов и форсированных вариантов пришлось бы добавить к формуле оценочной функции слишком много слагаемых, поэтому лучше учесть их путем обычного просмотра отдельных вариантов, т. е. сделать так, как делает игрок-шахматист, рассчитывая варианты. Нужно исследовать определенное число вариантов ход за ходом до тех пор, пока не получится сравнительно спокойная позиция, к которой затем применить некоторую оценочную функцию. Затем игрок выбирает вариант, который дает наибольшую оценочную функцию для него, в то время как противник старается уменьшить значение этой функции.

Этот процесс может быть описан математически. Сначала не будем учитывать того, что оценочная функция $f(P)$ может быть приложима только к спокойным позициям. Стратегия игры, основанная на $f(P)$ и действующая на один ход вперед, описывается следующим образом.

Пусть $M_1, M_2, M_3, \dots, M_s$ будут ходы, которые могут быть сделаны в позиции P и пусть $M_1P, M_2P, \dots$ и т. д. обозначают символически результирующие позиции, которые получаются из исходной ходами M_1, M_2 и т. д. соответственно. Будем выбирать ход M_m , для которого значение функции $f(M_mP)$ максимально.

Более глубокая стратегия должна принимать в расчет ответы противника. Пусть $M_{i1}, M_{i2}, \dots, M_{is}$ — возможные ответы черных, если белые выбрали ход M_i . Черные должны играть с таким расчетом, чтобы минимизировать $f(P)$. Кроме того, черные должны ходить после того, как белые сделают свой ход. Следовательно, если белые выбирают ход M_i , то они должны допустить возможность того, что черные выберут ход M_{ij} такой, что значение $f(M_{ij} M_iP)$ минимально. Белые должны выбрать свой первый ход с таким расчетом, чтобы это значение было максимальным после выбора черными лучшего ответа. Таким образом, белые должны выбрать ход M_i , который максимизирует величину

$$\min_{M_{ij}} f(M_{ij} M_iP).$$

Процесс, описанный выше, показан для простого случая на рис. 2. Точка слева обозначает исходную позицию. Предполагается, что в этой позиции имеется для белых три хода, указанных на рисунке тремя сплошными линиями; в каждой из получившихся позиций черные имеют по три возможности, обозначенные пунктирными линиями. Всем возможным позициям, которые могут получиться в результате хода белых и ответа черных, соответствует девять точек справа, а числа у этих точек обозначают величину оценочной функции, соответствующую этим позициям. Минимум по верхним трем точкам дает 0,1, что соответствует выбору белыми верхнего хода

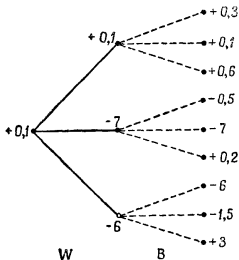
и затем лучшему ответу черных. Аналогично второй и третий ходы приводят к величинам -7 и -6 . Выбирая теперь максимум по ходам белых, находим, что это есть $0,1$, а, следовательно, белым надо выбрать первый ход.

Таким же образом стратегия, учитывающая два хода (просмотр всех вариантов на два хода вперед), задается так:

$$\max_{M_i} \min_{M_{ij}} \max_{M_{ijk}} \min_{M_{ijkl}} f(M_{ijkl} M_{ijk} M_{ij} M_i P) \quad (1)$$

Порядок следования операций максимизации и минимизации функции здесь важен. Это вытекает из того факта, что ходы белых и черных чередуются.

Машина, работающая по такой двухходовой стратегии, должна вычислить все варианты на два хода (с каждой стороны) и все



Р и с. 2.

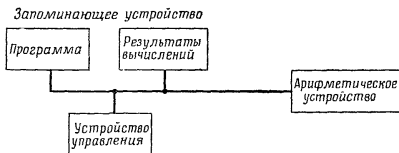
результатирующие позиции. Для каждой из этих позиций подсчитывается оценочная функция. Фиксируется все, кроме последнего хода черных, который варьируется, и выбирается ход, минимизирующий f . Это, согласно оценочной функции, самый сильный ход, который могут сделать черные в данном исследуемом варианте. Затем выбирается следующий второй ход за белых, и процесс повторяется для возможных вторых ходов черных. После перебора всех вторых ходов белых каждому из них ставится в соответствие определенное значение оценочной функции (то, которое соответствует лучшему следующему ответу черных). Далее, таким же процессом выбирается среди вторых белых ходов в каждом варианте тот, который дает максимум f (после выбора черными наилучшего второго хода в каждом случае). Продолжая таким образом, машина

идет обратно к начальной позиции и выбирает лучший первый ход за белых. Затем делает этот ход. Этот процесс очевидным образом обобщается на любое число ходов.

Такую стратегию, в которой все варианты рассматриваются на определенное число ходов вперед и затем ход определяется по формуле вида (1), будем называть стратегией типа А. Стратегия типа А имеет определенные слабые стороны, которые будут рассмотрены позднее, но она в принципе проста. Покажем, как составить программу для ее реализации.

5. Составление программы для универсальной вычислительной машины для реализации стратегии типа А

Пусть имеется универсальная цифровая машина, изображенная схематично на рис. 3 и обладающая свойствами.



Р и с. 3.

1) Имеется большая внутренняя память для хранения чисел. Она разделена на некоторое число ячеек, каждая из которых способна содержать, скажем, десятиразрядное число. Каждой ячейке присвоен адрес в памяти.

2) Машина имеет арифметическое устройство, которое может выполнять элементарные операции сложения, умножения и т. д.

3) Вычислительная машина работает, повинаясь программе. Программа состоит из последовательности элементарных команд. Обычная команда имеет вид:

Сл 372, 451, 133.

Это означает: выбрать содержимое ячеек 372 и 451, сложить эти числа и записать их сумму в ячейку 133.

Другой тип команд содержит в себе выбор, например,

Ср 291, 118, 345.

По этой команде машина сравнивает содержимое ячеек 291 и 118. Если первое число больше второго, машина выполняет следующую по порядку команду программы, если нет, то следующей выполняется команда из ячейки 345. Этот тип команд дает возможность машине сделать выбор из двух возможностей в зависимости от предыдущих результатов. Будем предполагать, что при помощи команд можно переносить числа из ячейки в ячейку, производить арифметические операции и осуществлять выбор (в только что указанном смысле).

Наша задача состоит в представлении шахматной игры в виде чисел и операций над числами и в записи стратегии в виде последовательности команд. Не приводя детали, опишем в общих чертах программы. Ясно, что конечная программа должна быть записана в виде команд.

Эта, так сказать, прокрустова задача представления шахмат в машине, предназначенной для выполнения математических вычислений, вызвана экономическими соображениями. Лучше всего было бы построить специальную вычислительную машину для шахмат, содержащую вместо арифметического устройства «шахматное устройство», специально сконструированное для выполнения простых шахматных вычислений. Несмотря на то что при этом, несомненно, будет достигнуто большое увеличение скорости выполнения операции, стоимость такой вычислительной машины будет так высока, что ее создание окажется невозможным. Предполагается все же провести ряд экспериментов на одной из цифровых вычислительных машин, которая сейчас конструируется.

Игра в шахматы может быть разделена на три фазы: дебют, середину и окончание. На разных фазах приложимы различные принципы игры. В дебюте, который обычно заканчивается в течение примерно десяти ходов, главной целью является развитие фигур для занятия хороших позиций. В середине игры превалируют тактические удары и комбинации. К концу этой фазы обычно большинство фигур разменивается, остаются только короли, пешки и, может быть, одна или две фигуры с каждой стороны. Эндшпиль в основном связан с пешечными превращениями, где становится важным точное определение таких возможностей, как цугцванг, пат и т. д.

Ясно, что при таком разнообразии стратегических задач на различных стадиях игры должны быть использованы различные программы. Будем в основном касаться середины игры и совсем не будем рассматривать окончания. Не видно, однако, причин, по которым не могла бы быть построена и достаточно хорошо запрограммирована стратегия окончания игры.

Поле на шахматной доске может быть занято тринадцатью различными способами: либо оно пусто (0), либо на нем стоит одна из шести возможных белых фигур ($p=1$, $K=2$, $C=3$, $L=4$, $\Phi=5$,

$Kp=6$), либо одна из шести черных фигур ($p=-1$, $K=-2$, $C=-3$, $L=-4$, $\Phi=-5$, $Kp=-6$). Таким образом, состояние поля можно охарактеризовать приписыванием ему целого числа от -6 до $+6$. 64 поля можно пронумеровать так, как показано на рис. 4. Положение всех фигур тогда задается в виде последовательности 64 чисел, каждое из которых находится между -6 и $+6$. Для такого представления достаточно всего 256 двоичных разрядов. Хотя это

Черные

70	71	72	73	74	75	76	77
60	61	62	63	64	65	66	67
50	51	52	53	54	55	56	57
40	41	42	43	44	45	46	47
30	31	32	33	34	35	36	37
20	21	22	23	24	25	26	27
10	11	12	13	14	15	16	17
00	01	02	03	04	05	06	07

Белые

Р и с. 4.

Обозначения фигур

символ	П	К	С	Л	Ф	Кр
белые	1	2	3	4	5	6
черные	-1	-2	-3	-4	-5	-6

0=пустая клетка

Обозначение ходов (исходное положение, результат хода, новая фигура, если пешка прошла в ферзи)

$e2-e4 \rightarrow (14, 34, -)$
 $e7-e8\Phi \rightarrow (64, 74, 5)$

кодирование не является оптимальным, оно достаточно удобно для вычислений. Еще одно число λ будет иметь значение $+1$, или -1 в зависимости от того, принадлежит ход белым или черным соответственно. Затем необходимо добавить данные, относящиеся к возможности рокировки (двигались ли белые или черные короли и ладьи), а также к взятиям на проходе (т. е. описание последнего

хода). Однако здесь эта информация опущена. Во введенных обозначениях начальная шахматная позиция имеет вид

4,	2,	3,	5,	6,	3,	2,	4;	1,	1,	1,	1,	1,	1,	1,	1;
0,	0,	0,	0,	0,	0,	0,	0;	0,	0,	0,	0,	0,	0,	0,	0;
0,	0,	0,	0,	0,	0,	0,	0;	0,	0,	0,	0,	0,	0,	0,	0;
-1,	-1,	-1,	-1,	-1,	-1,	-1,	-1;	-4,	-2,	-3,	-5,	-6,	-3,	-2,	-4;

+1 ($=\lambda$).

Любой ход (кроме рокировки и превращения пешки) можно описать указанием того, с какого поля и на какое переставляется фигура. Всего на доске 64 поля, следовательно, достаточно 6 двоичных разрядов для указания одного поля и всего 12 двоичных разрядов нужно для описания хода. Таким образом, первый ход $e2-e4$ записывается в виде 1,4; 3,4. Чтобы закодировать превращение пешки, необходимо добавить 3 двоичных разряда для указания фигуры, в которую она превращается. Рокировку можно определить по движению короля (это единственный случай, когда король может двинуться сразу на два шага). Всякий ход, следовательно, можно задать выражением (a, b, c) , где a и b обозначают поля, а c указывает фигуру в случае превращения пешки.

Полная программа для стратегии типа А состоит из девяти программ, которые обозначим $T_0, T_1, T_2, T_3, T_4, T_5, T_6, T_7, T_8$ и основной программы T_9 . Основные функции этих программ следующие.

T_0 — выполняет ход (a, b, c) в позиции P , определяя новую позицию.

T_1 — составляет список возможных ходов пешек с поля (x, y) в позиции P .

T_2-T_6 — составляют аналогичные списки для других фигур: коней, слонов, ладей, ферзей и короля.

T_7 — составляет список всех возможных ходов в данной позиции.

T_8 — вычисляет оценочную функцию $f(P)$ для данной позиции P .

T_9 — управляющая программа; выполняет вычисления, связанные с выбором максимумов и минимумов для определения искомого хода.

Имея данную позицию P и ход (a, b, c) во внутренней памяти, машина может найти следующую позицию, выполняя программу T_0 .

1) В памяти машины отыскивается поле a позиции P .

2) Число x в этом квадрате заменяется на 0 (пустой квадрат).

3) а) Если $x=1$ и первая координата a равна 6 (белая пешка достигает поля превращения) или если $x=-1$ и первая координата a равна 1 (черная пешка достигает поля превращения), то в поле b записывается число c (заменяя то число, которое было там записано).

б) Если $x=6$ и $a-b=2$ (короткая рокировка белых), то в поля 04 и 07 записывается 0, а в поля 06 и 05—числа 6 и 4 соответственно. Аналогично делается в случае $x=6$ и $b-a=2$ (длинная рокировка белых) и $x=-6$, $a-b=\pm 2$ (короткая или длинная рокировка черных).

в) Во всех остальных случаях в поле b записывается x .

4) Изменяется знак λ .

Для фигуры каждого типа имеется программа, определяющая ее возможные ходы. Типичным примером является программа T_3 для определения возможных ходов слона, которую кратко можно описать так. Пусть (x, y) координаты поля, на котором стоит слон.

1) Вычислим $(x+1, y+1)$ и посмотрим содержимое u этого поля в позиции Р.

2) Если $u=0$ (пустое поле), внесем в список ход (x, y) , $(x+1, y+1)$ и вернемся в начало, поставив $(x+2, y+2)$ вместо $(x+1, y+1)$. Если λu — положительно (наша фигура в поле), перейдем к 3). Если λu — отрицательно (в поле фигура противника), внесем в список ход и перейдем к 3). Если поля не существует, перейдем к 3).

3) Вычислим $(x+1, y-1)$ и проведем аналогичные рассуждения.

4) Аналогично с $(x-1, y+1)$.

5) Аналогично с $(x-1, y-1)$.

По такой программе составляется список возможных ходов слона в позиции Р. Подобные программы составят списки возможных ходов для любых других фигур. При этом имеются значительные возможности для упрощения этих программ, например программа T_5 для ферзя может быть комбинацией программ T_3 для слона и T_4 для ладьи.

Используя программы для фигур $T_1, \dots, T_6$ и логическую программу T_7 , машина может составить список всех возможных ходов в любой данной позиции Р. Управляющая программа T_7 коротко может быть описана так (опуская детали).

1) Обращаемся к полю (1,1) и определяем его содержимое x .

2) Если λx положительно, обращаемся к соответствующей программе T_x и по окончании ее работы возвращаемся к 1), прибавив 1 к номеру поля. Если λx есть нуль или отрицательно, возвращаемся к 1), прибавив 1 к номеру поля.

3) Проверяем, допустим ли каждый из выписанных ходов, и выбрасываем недопустимые ходы. Это делается путем выполнения каждого хода в позиции Р (по программе T_0) и проверкой, находится ли король под шахом или нет.

Машина может играть в шахматы, пользуясь только программами $T_0, \dots, T_7$ (т. е. соблюдая правила игры и делая каждый раз просто случайно выбранный допустимый ход). Уровень игры по

такой стратегии очень низок¹⁾. Автор сыграл несколько игр против такой случайной стратегии и обычно делал противнику мат в четыре-пять ходов. Следующая партия иллюстрирует полную бессмысленность игры по такой стратегии:

Белые (ход выбирается случайным образом)		Черные
1. g	3	e 5
2. d	3	C c5
3. C	d2	Ф f6
4. K	c3	Ф: f2 ×

Теперь вернемся к стратегии, основанной на оценочной функции $f(P)$. Программа T_8 производит оценку позиции согласно значению $f(P)$. Очевидно, это можно сделать, просматривая поля и суммируя указанные в оценочной функции члены. Нетрудно учесть такие факторы, как сдвоенные пешки и т. д.

Последняя управляющая программа T_9 необходима для того, чтобы выбрать ход в соответствии с минимаксным процессом, описанным выше. В случае стратегии, учитывающей по одному ходу с каждой стороны, T_9 работает следующим образом:

1) Составляет список допустимых ходов (используя T_7) в данной позиции.

2) Делает первый ход в списке по программе T_0 , получая позицию M_1P .

3) Составляет список ходов черных в позиции M_1P .

4) Делает первый ход за черных, получает позицию $M_{11}M_1P$ и оценивает ее по T_8 .

5) Делает второй ход за черных и оценивает полученную позицию.

6) Сравнивает оценки и выбрасывает ход с меньшей оценкой (т. е. остается лучшая оценка за черных).

7) Выполняет для третьего хода сравнение с оставшейся величиной и т. д.

8) Когда ходы черных исчерпаны, будет оставлен один ход с его оценкой. Процесс затем повторяется для второго хода белых.

9) Конечные оценки этих двух вычислений сравниваются, и оставляется тот из ходов белых (вместе с соответствующей оценкой), у которого оценочная функция больше.

10) Этот процесс повторяется для всех ходов белых и выбирается лучший (он остается после перебора всех возможных ходов белых). Это и есть тот ход, который надо сделать.

¹⁾ Имеется вероятность порядка 10^{-75} , что при такой игре можно выиграть у Ботвинника, но возможно играть еще хуже, руководствуясь стратегией, которая помогает противнику. Например, стратегия белых в следующей партии:

1. f3 e5. 2. g4 Фh 4 ×.

Эти программы в высокой степени цикличны, следовательно, если они хорошо составлены, не должны занимать много места в памяти машины.

Можно оценить объем памяти для записи позиций и промежуточных результатов, когда просматриваются все варианты на три хода вперед (с каждой стороны). Вероятно, должны запоминаться три позиции: начальная позиция, последняя позиция (которая теперь оценивается) и предшествующая ей. Это потребует около 800 двоичных разрядов. Кроме того, нужно помнить пять списков ходов, каждый из которых требует около $30 \times 12 = 360$ двоичных разрядов, а все они потребуют 1800 двоичных разрядов. Наконец, чтобы сделать выбор и оценку в необходимых вычислениях, требуется около 200 двоичных разрядов. Значит, всего потребуется порядка 3000 двоичных разрядов.

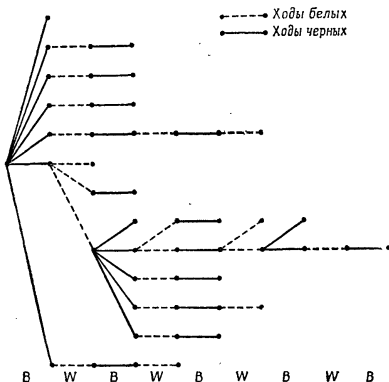
6. Улучшение стратегии

К сожалению, машина, действующая по стратегии типа А, будет и слабым и медленным игроком. Она будет играть медленно, так как даже если на оценку одной позиции она затратит одну микросекунду (что очень оптимистично), то для оценки 10^9 позиций, которые надо сделать при трехходовом (за каждую сторону) варианте стратегии, ей потребуется более 16 минут на ход, т. е. около 10 часов для «обдумывания» ходов за свою сторону на 40-ходовую партию.

Машина будет малоискусным игроком, потому что она будет просматривать варианты только на три хода вперед. Кроме того, наша стратегия не дает ей возможности выделить спокойные позиции, в которых разумно производить оценку. Машина действует крайне неэкономным способом по такой программе — она считает все варианты ровно на три хода вперед и затем прекращает расчеты (даже если противник под шахом). Хороший шахматист проверяет только несколько избранных вариантов и просчитывает их на разумную глубину. Чемпион мира может рассчитать комбинации на глубину, скажем, 15—20 ходов. Некоторые варианты, данные Алехиным («Мои самые лучшие шахматные партии 1924—1937 гг.»), имеют такую длину. Конечно, на такую глубину исследуется только небольшое количество вариантов. В играх любителей редко проверяются варианты более чем на 6—8 ходов, и то только тогда, когда ходы носят форсированный характер (ответы противника очень ограничены).

Вообще говоря, когда имеется немного угроз и форсированных ходов, большинство исследуемых вариантов имеет глубину один-два хода, и только около полудюжины форсированных вариантов рассчитываются на 4—5 ходов.

С этой точки зрения интересно признание Р. Файна¹⁾, ведущего американского шахматиста: «Очень многие полагают, что шахматные мастера предусматривают все или почти все, что, играя h3 на тринадцатом ходу, они предвидят, что этот ход будет им полезен, чтобы сохранить отход королю через двадцать ходов. Некоторые



Р и с. 5.

думают, что когда они играют 1. e4, то делают это для того, чтобы предотвратить ход Kd5 при двенадцатом ходе черных, или что все вычислено до того момента, когда ладейная пешка ферзевого фланга станет ферзем на один ход раньше, чем коневая пешка королевского фланга противника. Все это, конечно, чистая фантазия. Лучше всего разобрать основные варианты на пару ходов, но попытаться рассчитать форсированные варианты до конца».

Количество вариантов, отбираемых шахматными мастерами для исследования в различных позициях, было изучено экспериментально де Гроотом²⁾. Он показывал различные типовые позиции шахматистам и просил их выбрать лучший ход, продельвая вслух анализ позиции, который они проводят про себя.

¹⁾ Fine R., Chess the easy way, 1942, David McKay.

²⁾ De Groot A., Het denken van den Schaker, Amsterdam, 1946, 207.

Таким способом может быть установлено число и глубина исследуемых вариантов. Рис. 5 показывает результаты одного из таких экспериментов. В этом случае шахматист проверил 16 вариантов, а глубина анализа колебалась от $1\frac{1}{2}$ (один ход черных) до $4\frac{1}{2}$ (пять ходов черных и четыре белых). Всего рассматривалось 44 позиции.

Из этих замечаний следует, что для увеличения скорости и силы игры машина должна:

1) возможно глубже проверять форсированные варианты и оценивать только такие позиции, где установлена некоторая квазистабильность;

2) выбирать варианты для исследования согласно некоторому алгоритму, который позволит машине не расходовать напрасно время на исследование бесперспективных возможностей.

Стратегию с этими двумя усовершенствованиями будем называть стратегией типа В. Нетрудно составить программу, реализующую стратегию такого типа. Сначала определим функцию $g(P)$, которая определяет, стабильна ли в какой-то мере данная позиция (нет фигур под ударом и т. д.).

Грубым приближением такой функции может служить функция

$$g(P) = \begin{cases} 1, & \text{если какая-либо фигура атакована фигурой противника} \\ & \text{меньшей цены, или большее количество фигур атакует} \\ & \text{какую-либо фигуру, чем защищает, или король под} \\ & \text{шахом.} \\ 0 & \text{в противном случае.} \end{cases}$$

Используя эту функцию, можно исследовать варианты до тех пор, пока $g(P)$ не станет равной нулю, продвигаясь, однако, не меньше чем на два хода, и не больше, скажем, чем на 10.

Следующее усовершенствование требует введения функции $h(P, M)$, чтобы решить, нужно ли в позиции P исследовать ход M . Важно, чтобы при оценке не выбрасывались ходы, которые выглядят плохими только на первый взгляд, например ход, подставляющий под удар фигуру; часто такие ходы бывают в действительности очень сильными, так как подставленная фигура не может быть взята безнаказанно.

«Всегда давайте шах, возможно это мат» — вот иронический совет, даваемый начинающим, использующий их предрасположение к бесполезным шахам. «Всегда исследуйте шах, он может привести к мату» — разумный совет для любого игрока. Шах придает течению партии форсированный характер. Ответы противника крайне ограничены, например, он никогда не может ответить контратакой. Это означает, что вариант, начинающийся с шаха, легче рассчитать, чем другие. Аналогично взятия фигур и атаки на фигуру, имеющую большую цену, угрозы мата и т. д. ограничивают ответы противника и должны быть исследованы независимо от того, выглядит ход хорошим или нет. Следовательно, функция $h(P, M)$ должна

давать большую величину для всех форсированных ходов (шахи, взятия и атакующие ходы), для развивающих ходов, среднюю величину для защитных ходов и низкую для других. При исследовании вариантов машина должна подсчитывать функцию $h(P, M)$ и использовать ее для выбора тех вариантов, которые надо проверить. При продвижении в глубь варианта требования к функции h увеличиваются, что приводит к тому, что исследуется все меньшее число вариантов. Таким образом, машина, начиная с рассмотрения каждого своего хода, рассматривает только форсированные ответы противника и т. д. Такой процесс сильно увеличит эффективность вычислений.

Можно считать, что универсальная вычислительная машина с учетом этих двух дополнений к программе будет играть довольно сильно и со скоростью, сравнимой с человеческой. Нужно отметить, что машина при этом будет иметь следующие преимущества по сравнению с человеком.

1. Огромная скорость выполнения отдельных вычислений.
2. Свобода от ошибок. Единственными ошибками являются те, которые были допущены в программах, тогда как шахматисты постоянно допускают простые и очевидные просмотры.
3. Отсутствие лени. Шахматисту очень легко сделать инстинктивно ход без надлежащего анализа позиции.

4. Отсутствие «нервов». Шахматисты склонны переоценивать свои шансы в «выигрышной» позиции или теряться в «проигрышной».

Эти качества должны в какой-то мере компенсировать отсутствие у машины той гибкости, воображения, логического мышления и возможности обучения, которые присущи человеку. Вообще говоря, человек, составивший программу, может вычислить ход, который машина должна сделать в любой позиции, и, следовательно, в некотором смысле он может играть так же хорошо. В действительности, однако, вычисления могут оказаться практически невыполнимыми, так как потребуют много времени. При правильном способе сравнения, когда машине и автору программы дается равное время на выбор хода, она может сыграть гораздо лучше.

7. Изменения в игре и в стиле

Если машина играет по программе, описанной выше, она будет всегда выбирать один и тот же ход в одной и той же позиции. Если противник будет повторять свои ходы, это приведет к повторению партии. Желательно, чтобы этого не получилось, так как противник, победив один раз, будет выбирать тот же вариант и неизменно добиваться победы в результате того, что в некоторой частной позиции машина выбирает очень слабый ход.

Одним из путей предотвращения такого поведения машины является введение случайных элементов в ее программу. Всякий раз, когда имеется два или более ходов, которые имеют приблизительно равную оценку, согласно вычислениям машины, она делает выбор случайным образом. При рассмотрении позиции второй раз машина может выбрать другой ход.

Начало игры — другое место, в котором можно вводить случайные методы выбора варианта. Представляется желательным иметь некоторое количество стандартных вариантов в дебютах, которые записаны в медленной памяти машины. Вероятно, достаточно иметь несколько сотен таких вариантов. Первые ходы (до тех пор, пока противник не отклонится от «книги» или не исчерпается записанный вариант) машина делает по памяти. Это совсем не «жульничество», так как шахматисты поступают таким же образом, разыгрывая дебют.

Интересно, что «стиль» игры машины может быть очень легко изменен варьированием некоторых коэффициентов и численных значений логических факторов, включенных в оценочную функцию и в другие программы. Придавая большое значение слабостям позиций, можно заставить машину тяготеть к позиционной игре. Заставляя ее более интенсивно изучать форсированные варианты, лучшим ее комбинационную игру. Далее, сила игры может быть легко отрегулирована изменением глубины вычислений и опусканием или введением членов в оценочную функцию.

Наконец, надо отметить, что машина такого типа будет играть «блестяще» в определенных пределах. Она будет готова пожертвовать ферзя или другую фигуру для достижения позднее материального преимущества или будет объявлять мат, лишь бы только окончание комбинации лежало в сфере ее вычислительных возможностей. Основной слабостью такой машины будет являться отсутствие учета сделанных ошибок. Единственным путем улучшения ее игры является изменение программы. Высказываются мысли о создании программы, которая самоулучшается, но, хотя это кажется возможным, предложенные на сегодня методы не представляются практичными. Одно из предложений заключается в создании программы высшего уровня, которая изменяет слагаемые и коэффициенты, включенные в оценочную функцию в зависимости от результатов игры машины. Для достижения наибольшего процента побед могут быть введены лишь небольшие изменения в заранее выбранные члены и величины.

8. Другой тип стратегии

Описанные выше стратегии не исчерпывают, конечно, всех возможностей. Несомненно существуют другие типы стратегий, при которых машинное время будет использоваться более эффек-

тивно. Даже с теми улучшениями, которые были рассмотрены выше, стратегия машины носит медлительный характер. Слишком много «грубой силы» тратит она на вычисления, а не на логический анализ позиции. Она играет подобно начинающему шахматисту, которому рассказали о принципах игры и который владеет огромной энергией и точностью расчетов, но не имеет опыта игры. Шахматный мастер, с другой стороны, знает сотни, а может быть, и тысячи стандартных позиций, привычных комбинаций и типовых маневров, которые встречаются неоднократно в партиях. Имеются, например, стандартные жертвы коня на f7 или слона на h7, стандартные маты, например мат Филидора, маневры, связанные с вилками, связками, вскрытиями, превращениями и т. д. В данной позиции он обнаруживает много сходства со знакомыми ему случаями и это направляет его мысль на исследование тех вариантов, в которых вероятность успеха наибольшая.

Нет никаких причин, препятствующих написанию программы, реализующей стратегию, основанную на использовании «типовых позиций». Однако на этом пути потребуются тщательный анализ шахматной игры. Хотя существует много книг, анализирующих комбинационную игру и середину игры, они написаны для восприятия человеком, а не машиной. Вполне возможно указать человеку один или два специфических примера общей ситуации, чтобы дать ему понять и научить его применять общий принцип, заложенный в них. Для вычислительной машины должна быть дана точная и совершенно полная информация со всеми ограничениями, специальными случаями и т. д., которые надо принять в расчет. Мы склонны верить, однако, что программа с использованием такой информации сильно увеличила бы эффективность игры.

Для программирования такой стратегии можно полагать, что любая позиция сопровождается соответствующим образом закодированным тщательным анализом ее тактической структуры. Такие аналитические данные утверждают, например, что черный конь на f6 связан ферзем, что белая ладья на e1 не может оставить своей горизонтали вследствие возможности мата на f8, что белый конь на h4 не имеет возможности пойти и т. д. Короче, описываются все факты, которым игрок уделяет внимание при анализе тактических возможностей. Эти данные используются программой и изменяются ею по ходу игры. Они служат для вызова различных программ в зависимости от природы рассматриваемой позиции. Связанную фигуру нужно атаковать. Если ладья охраняет первую горизонталь, она не может защищать пешку по вертикали и т. д. Вот такими «рассуждениями» машина может находить ходы, которые нужно анализировать.

Но из сказанного не следует, что надо стараться провести в программе стратегию, моделирующую наше собственное поведе-

ние при выборе хода за доской. Прежде всего, необходимо учитывать сильные и слабые стороны вычислительной машины. Она сильна огромной скоростью, с которой производятся отдельные операции, высокой точностью вычислений, но имеет слабые аналитические способности и возможности отождествления отдельных элементов позиции. Следовательно, она должна выполнять больше грубых вычислений, чем человек, при этом она может просмотреть до 10^8 вариантов при каждом ходе. Небольшое добавление избирательности в игре явится хорошим средством избежания грубых зевков и ошибок.

ПРИЛОЖЕНИЕ

ОЦЕНОЧНАЯ ФУНКЦИЯ ДЛЯ ШАХМАТ

Оценочная функция $f(P)$ должна принимать в расчет такие преимущества и недостатки позиций, которые могут сказаться через большее количество ходов, чем то, на которое проводится просмотр вариантов. Следовательно, она прежде всего имеет дело с позиционной, стратегической оценкой позиции, а не с тактической комбинационной борьбой. При этом между ними нельзя провести четкую грань, так как они часто совпадают. Ясно, однако, что в функции $f(P)$ могут учитываться следующие свойства позиции.

1. Материальное преимущество.
2. Пешечная структура:
 - а) отставшие, изолированные и сдвоенные пешки;
 - б) относительный контроль центра (пешки на e4, d4 и c4)¹⁾;
 - с) слабость пешек вблизи короля (т. е. продвинутые вперед слоновые, коневые и ладейные пешки);
 - д) пешки на полях, цвет которых отличен от цвета слона;
 - е) проходные пешки.
3. Положение фигур:
 - а) продвинутый конь (на e5, d5, c5, f5, e6, d6, c6, f6) особенно такой, которого нельзя атаковать неприятельской пешкой и который защищен своей пешкой;
 - б) ладья на открытой или полуоткрытой линиях;
 - с) ладья на седьмой горизонтали;
 - д) сдвоенные ладьи.
4. Защита, атака и право выбора:
 - а) фигуры, которые необходимы для защитных функций и, следовательно, связанные и ограниченные в передвижении;

¹⁾ Здесь и всюду ниже при указании конкретных полей или линий имеются в виду белые фигуры. Для черных фигур подразумеваются симметричные поля и линии.— *Прим. перев.*

б) атака на фигуры, которая дает игроку¹ возможность выбора размена;

с) атака полей рядом с королем;

д) связи; здесь подразумеваются скрывающие связи, когда связанная фигура имеет цену не большую, чем связывающая. Например, конь, связанный слоном.

5. Мобильность.

Все эти факторы приложимы к средней стадии игры, а в начале игры и в окончаниях разумно использовать другие принципы. Относительные веса, которые необходимо придать перечисленным выше качествам, — это вопрос, требующий доработки; он может быть решен экспериментальным путем. Имеется еще много других факторов, которые было бы интересно учесть. Наиболее грозное тактическое оружие—вскрытые шахи, вилки и связи фигурой низшей цены—здесь опущено, так как его лучше всего учесть исследованием конкретных вариантов.

ЛИТЕРАТУРА

- Chernev, Curious chess facts, The black Knight Press, 1937.
De Groot A. D., Het Denken van den Schaker, Amsterdam, 1946a, 17—18; там же, 1946b, 207.
Fine R., Chess the easy way, David Mc Kay, 1942, 79.
Hardy G., Wright E., The theory of numbers, 1938, Oxford, 116.
von Neuman J., Morgenstern O., Theory of games and economic behavior, Princeton, 1944, 125.
Vigneron H., Les automates, La Nature, 1914.
Wiener N., Cybernetics, John Wiley, 1948; русский перевод: Винер Н., Кибернетика, Советское Радио, М, 1958.

Конструирование играющих машин на первый взгляд может показаться скорее интересным времяпрепровождением, чем серьезной научной задачей, и действительно, для многих ученых — любителей и профессионалов — этот увлекательный предмет стал любимым занятием. Однако эта работа имеет свою серьезную сторону и важную цель, и по крайней мере четыре или пять университетов и лабораторий работают сейчас в этом направлении. Если бы был жив Бенджамин Франклин, то, я уверен, он заинтересовался бы этой проблемой, так как в ней сливаются два направления, которым он отдавал энергию и время. Всем нам известны его достижения как ученого и изобретателя, но не столь широко известно, что он был также и сильным шахматистом. Он является автором занимательного очерка под названием «Мораль в шахматах», представляющего собой смесь шахматной теории и дипломатии, которому вполне подошел бы подзаголовок «Как можно быть счастливым, даже будучи шахматистом».

Одним из наиболее важных технических достижений за последние двадцать лет является развитие универсальных электронных вычислительных машин. Эти машины могут автоматически выполнять длинные последовательности вычислительных операций со скоростью тысяч операций в секунду. Последовательность команд, сообщаемых машине, что именно ей следует делать, называется программой. Если машине предстоит решить задачу, то сначала нужно составить программу, переводящую процесс нахождения решения в ряд простых операций. Эти основные операции могут включать элементарные арифметические действия — сложение, умножение и т. п., а также операции «передачи управления», позволяющие машине выбрать одну из двух возможностей в зависимости от результатов предшествующих вычислений. Когда программа

¹⁾ Shannon C., Game Playing Machines, *Journal of the Franklin Inst.*, 260, № 6 (1955), 447. Эта лекция была прочитана 19 октября 1955 г. в связи с вручением автору медали Стюарта Баллантина.

введена в машину, машина с большой скоростью выполняет команды одну за другой.

Электронные вычислительные машины обычно используются для решения численных задач, возникающих в науке или в промышленности. Однако основная схема этих машин является настолько гибкой и универсальной по своей идее, что на них можно программировать многие задачи, совершенно не связанные с вычислениями, такие, как перевод с одного языка на другой, анализ логической ситуации или решение игровых задач. Те же самые команды, которые применяются для записи программ и решения численных задач, могут символизировать операции над абстрактными объектами, такими, как слова языка или позиции на шахматной доске.

Составление программ для решения таких неарифметических задач на вычислительных машинах важно по целому ряду причин. Оно расширяет наши знания относительно возможностей того поразительно гибкого орудия, каким является универсальная вычислительная машина; представляется несомненным, что мы лишь поверхностно знакомы с возможностями таких вычислительных устройств и каждая новая сфера их применения ведет к расширению наших знаний. Кроме того, это более широкое использование вычислительных машин приводит к полезным изменениям в их конструкции, порождает новые типы операций, которые могут использоваться в необычных программах и даже в обычных численных задачах. Наконец, можно надеяться, что исследования в области конструирования играющих машин могут привести к углублению наших знаний о работе человеческого мозга. Разумеется, было бы наивным ожидать, что мозг действует так же, как машина, предназначенная для ведения игр или сколько-нибудь аналогично ей. Тем не менее несомненно, что создание любой обучающейся машины будет освещать путь к пониманию работы мозга.

Пожалуй, самой первой играющей машиной был Маельзельский автомат для игры в шахматы. Это было устройство, сконструированное в 1769 году австрийцем фон Кемпеленом и широко демонстрировавшееся в Европе и Америке предпринимателем Маельзелем.

Большая механическая фигура, сидевшая за столом, играла с людьми в шахматы. Перед демонстрацией стол и фигура открывались, чтобы показать, что внутри никого нет. Машина обычно выигрывала партии против человека-шахматиста. В то время автомат вызвал подлинную сенсацию и предлагался ряд теорий с объяснением его работы. Среди них был, например, очерк Эдгара Аллана По, который вывел верное заключение (хотя отчасти при помощи неверных аргументов), что машина была мистификацией и в действительности управлялась человеком-шахматистом, искусно спрятанным внутри. Так оно на самом деле и было, а требуемое впечатление, как это бывает во многих трюках иллюзионистов, создава-

лось перемещением шахматиста внутри машины, в то время как различные отделения и двери открывались для обозрения. После многих приключений, в том числе нескольких игр с императором Наполеоном, автомат был помещен в Китайский музей в Филадельфии и, наконец, погиб при пожаре в 1854 году.

Несколько лет тому назад мое внимание привлек современный двойник Маэльзельского автомата. Один мой друг из Калифорнии написал мне, что там на местной выставке и по телевидению показывали машину, играющую в шашки. Победить ее было почти невозможно — даже игру с чемпионом США она свела вничью, — и большинство считало, что это в самом деле не что иное, как электронная вычислительная машина. Исследуя проблему программирования игры в шахматы и шашки, я был настроен довольно скептически, особенно из-за того, что машина играла так хорошо и была такой портативной. Поэтому я предложил произвести осмотр устройства. После изрядной «сыскной» работы мой друг, наконец, выследил игрока в шашки на старом складе. Он сообщил, что единственной «электронной» частью машины был электрический вентилятор для скрытого человека-оператора.

Если не принимать во внимание подобных мистификаций, то играющие машины в зависимости от степени их сложности можно подразделить на три основных типа. Самый простой тип представляют машины, предназначенные для игры в полностью проанализированные игры. Под этим подразумевается, что для сравнительно простой игры полностью известна стратегия, которая диктует соответствующий ход в каждой ситуации, возникающей в игре. Примерами таких игр являются «крестики и нолики» (tic-tac-toe), настольная игра Ним и ряд других математических игр. В подобных случаях можно представить известную стратегию в виде программы для универсальной или специализированной вычислительной машины так, чтобы машина делала правильный ход в любой момент игры. Этот тип машины, как правило, будет играть наилучшим образом. Машина будет выигрывать всякий раз, когда возможен выигрыш.

Одна из первых машин этого типа была создана приблизительно в 1914 г. испанским изобретателем Торресом-и-Квеведо. Эта машина проводила шахматный эндшпиль — король и ладья против короля. Этот эндшпиль сравнительно прост, и надлежащие ходы могут быть описаны несколькими простыми правилами. Машина Торреса переводила эти правила в переключательную схему. Хотя по современным стандартам это устройство выглядит простым, для своего времени это было, несомненно, выдающимся изобретением.

Другой игрой, поддающейся полному математическому анализу, является игра Ним. Было сконструировано много машин, играющих в эту игру. Как я полагаю, первая машина такого рода экспони-

ровалась на всемирной выставке в Нью-Йорке в 1939 г. Игру Ним математики исследовали давно, и оказалось, что правильную стратегию этой игры очень легко выразить в двоичной системе счисления. Релейные схемы естественнее всего могут быть приспособлены к двоичной системе, и поэтому было легко перевести математическую стратегию игры в Ним и в релейную схему. Большинство машин для игры в Ним играют наилучшим образом в том смысле, что они выигрывают, когда выигрыш вообще возможен, но они обычно представляют сделать первый ход человеку, так что последний может выиграть, если он играет безошибочно. Если противник делает хотя бы одну ошибку, машина перехватывает инициативу и выигрывает.

Любимой игрой всех, кто конструирует машины, является игра «крестики и нолики». Игру можно полностью проанализировать путем перечисления всех возможных партий. Если принять во внимание симметрию доски, то число таких вариантов сравнительно невелико. Одна из первых машин, играющих в крестики и нолики, была создана лет пятнадцать назад У. Кейстером, и здесь у меня имеется такая машина, реализующая электрическую схему Кейстера. При игре с машиной первый ход делает человек и машина всегда сводит игру по меньшей мере к ничейному результату. Если, однако, игрок допускает серьезную ошибку, то машина выигрывает. В машине имеется также «противообманное» устройство. Если попробовать сделать два хода подряд, то загорается лампочка «протеста».

Второй основной класс машин относится к играм, полный анализ которых неизвестен, но для которых известны общие принципы игры или принципы построения правильной стратегии. К таким играм относится большинство общеизвестных игр, как, например, шашки, шахматы, бридж и покер. Такая машина, играющая в шахматы, исследует в каждой данной позиции различные ходы, которые она может сделать, различные ответные ходы своего партнера и т. д. на два или три хода вперед. В конце каждой из таких проб машина может применить к результатам «оценочную функцию». В качестве своего хода она выбирает тот, который ведет к позиции с наивысшей оценкой, а относительно противника предполагается, что он играет так, чтобы свести эту оценку к минимуму. Это — известный в теории игр процесс минимакса. Поскольку оценочная функция или общие принципы игры не могут быть здесь непогрешимыми, сконструированная на их основе машина не будет играть безукоризненно. Однако можно ожидать, что она будет делать достаточно хорошие ходы, если общие принципы тщательно продуманы.

Примером системы, реализующей игру по общим стратегическим принципам этого типа, является программа игры в шашки, разра-

ботанная Ч. С. Стрэчи для использования на большой вычислительной машине. Первая игра, проведенная с помощью программы Стрэчи, показала, что игра может быть достаточно хорошей в дебюте и в миттельшпиле, но в эндшпиле очень плоха. Очевидно, что для заключительных этапов игры нужен существенно иной тип программ.

Другая программа для игры в шашки была разработана А. Л. Самуэлем, и имеются сведения, что между машинами Стрэчи и Самуэля организуется матч. О программе Самуэля рассказывают одну, может быть неправдоподобную, историю. Когда он впервые ввел свою программу в машину и нажал пусковую кнопку, чтобы машина сделала свой первый ход, она бешено проработала в течение нескольких минут, а потом напечатала ответ; «Сдаюсь»!

Другая играющая машина, принципиально совершенно иного типа была построена Э. Ф. Муром и мной. Это специализированная машина для игры в «Гекс». Гекс — это игра на доске, разбитой на правильные шестиугольники. Двое игроков по очереди передвигают фигурки людей по шестиугольникам. Один играет желтыми фигурами, а другой — синими. Цель желтых — образовать непрерывную цепь фигур от верхнего края доски до нижнего. Цель синих — образовать непрерывную цепь фигур от правого края доски до левого. После изучения этой игры у Мура и у меня возникла мысль представить доску в виде сети сопротивлений, а положение фигур — напряжением, подаваемым в соответствующие точки сети: положительным для желтых фигур и отрицательным для синих. Мы предполагали, что некоторые седловые точки в поле напряжения, образовавшемся в сети сопротивлений, будут соответствовать хорошим ходам в игре. Было построено простое устройство, реализующее эту сеть, и оказалось, что машина играла достаточно хорошо. Она выигрывала около 70% игр против посетителей лаборатории, если делала первый ход. Если машина делала второй ход, то выигрыш составлял 50% или меньше. Часто нас приятно удивляло, когда машина выбирала ходы, представлявшие на первый взгляд слабыми, но после тщательного анализа оказывавшиеся правильными и сильными. В одной из начальных позиций машина даже «открыла» такой ход, который был лучше всех ходов, применявшихся нами, и которым теперь обычно пользуются.

Третий и самый сложный тип играющей машины — машина, вырабатывающая свои собственные принципы игры. В программу включаются только правила игры и желаемый результат, а также некоторые общие принципы улучшения игры в результате учета полученного опыта. Машине затем предоставляется возможность играть много партий и предполагается, что методом испытаний и ошибок, путем подражания партнеру, путем анализа своих про-

игрышей и т. п. машина постепенно улучшит свое искусство игры. Хотя эта проблема обсуждалась много раз, машин, полностью подходящих под это определение, построено не было. Эта задача очень трудна, и стоимость проведения соответствующей программы исследований очень велика, но уже были построены несколько более элементарные обучающиеся машины. Рассмотрим две из них.

Д. У. Хегельбергер разработал интересное устройство, которое играет с человеком в «монетку»¹⁾. В обычной игре в монетку два игрока показывают монеты одновременно. Они могут показать монету либо гербом, либо решеткой. При совпадении выигрывает первый игрок, в противном случае — второй.

Получается своего рода психологическая игра в отгадывание, и выигрывает в конечном итоге тот, кто лучше может предсказать реакции противника. В машине Хегельбергера монеты заменены лампочками и переключателями. В машине есть регистры памяти, в которых запоминаются некоторые результаты игры с противником. Во время игры эти результаты анализируются машиной с целью отыскания некоторой системы или психологической тенденции у игрока — человека. Например, у одного может быть тенденция перейти от герба к решетке, когда он выиграл два раза подряд. Другой может делать обратное. Угадывающая машина ищет такие тенденции у противника, и те, что она находит, используются в последующей игре. Таким образом, она предполагает, что в будущем человек будет следовать тем тенденциям, которым он следовал в прошлом, и машина играет так, чтобы выиграть, если это повторение на самом деле произойдет. В результате игры с самыми различными людьми машина выиграла около 55% партий, что гораздо лучше, чем 50%, если бы выигрыш был чисто случайным.

Глубоко заинтересованный машиной Хегельбергера, я сконструировал другую машину для игры в монетку, использующую те же основные принципы, но гораздо более простую по объему памяти и другим деталям. После долгих споров о том, какая машина выигрывает, мы решили поставить эксперимент. Была сконструирована третья машина — посредник, которая могла передавать информацию между обеими машинами, считать очки и следить, чтобы игра велась по правилам. Все три машины были соединены вместе и играли несколько часов, причем зрители заключали небольшие пари и сопровождали игру громкими криками. Оказалось, что меньшая и, как предполагалось, менее «умная» машина победила большую со счетом около 55 на 45. Возможно, это произошло потому, что она быстрее меняла свои оценки действий партнера. Обе машины стараются найти схему игры противника, и как только одна находит

¹⁾ Хегельбергер Д. У., СИИР—автомат, экстраполирующий последовательности, Кибернетический сб., вып. 1, ИЛ, М., 1960, 275.—Прим. ред.

эту схему, другая начинает проигрывать и меняет свою стратегию. Поэтому более гибкий тип имеет определенное преимущество.

Машина для разгадывания лабиринта¹⁾ является другим примером обучающейся машины. Хотя она и не учится играть в какую-нибудь игру, она опытным путем учится решать определенный тип задач. Классический пример психологической задачи обучения состоит в том, что животное помещают в лабиринт и наблюдают за тем, какое время ему понадобится, чтобы найти дорогу к кормушке. Попыткой выяснения того, как подобную задачу сможет решить релейная схема, является машина для разгадывания лабиринта. Перегородки в лабиринте можно произвольно менять местами и получить около тысячи миллиардов различных лабиринтов, которые решает машина. Во время своего первого путешествия мышь следует принципам стратегии исследования, включающей большое количество испытаний и ошибок и неверных ходов, приводящих в тупики. Наконец, она достигает медного диска, представляющего кормушку. Если теперь поместить мышь в первоначальное положение, то она пойдет прямо к цели, без ошибочных ходов. Это показывает, что релейная схема запомнила верный путь. Более того, если мышь поместить в другую часть лабиринта, которую она исследовала раньше, она придет также прямо к цели. Если ее поместить в ту часть лабиринта, в которую она не попадала ранее, то она будет бродить до тех пор, пока не попадет на знакомую клетку, а оттуда пойдет прямо к цели. Если теперь поместить ее снова в исходную точку, она прямо идет к цели. Это показывает, что мышь добавила информацию об этой части лабиринта к предыдущему опыту. Наконец, если изменить лабиринт, мышь сначала пробует прежний путь, и когда ее постигает неудача, начинает вновь свои исследования и ищет другой путь к цели. Если кормушка закрыта со всех сторон и пути к ней нет, то мышь обследует все части лабиринта, до которых можно добраться.

Хотя развитие играющих машин в прошлом, в особенности за последние два десятилетия, представляется интересным и стимулирующим, думается, что лет через десять-двадцать эти устройства, несомненно, покажутся примитивными. Широкое использование универсальных вычислительных машин и быстрое развитие теории программирования обязательно приведут к важным достижениям в области создания машин, использующих общие принципы стратегии, и машин, обучающихся различным играм.

¹⁾ См. наст. сборник, стр. 223.—Прим. ред.

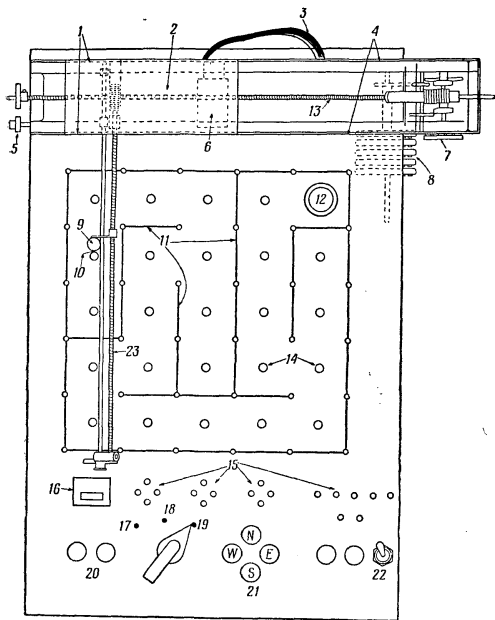
СООБЩЕНИЕ О МАШИНЕ, РЕШАЮЩЕЙ ЛАБИРИНТНУЮ ЗАДАЧУ¹⁾

Лабиринтная машина способна найти путь в лабиринте, применяя метод испытаний и ошибок. Машина запоминает решение, а в случае, когда ситуация меняется и решение оказывается неприемлемым, забывает его. Я думаю, что эта машина может представлять интерес с точки зрения связи ее с задачей обучения (и забывания) систем, содержащих обратные связи и работающих по методу испытаний и ошибок.

Как показано на рис. 1, на верхней панели машины расположен лабиринт, содержащий 5×5 квадратов. Лабиринт можно изменять любым желаемым образом, переставляя перегородки между двадцатью пятью квадратами. В лабиринте помещен чувствительный щуп, который может обнаруживать перегородки, когда он к ним подходит. Этот щуп связан с двумя двигателями, перемещающими его в двух взаимно перпендикулярных направлениях. Условимся считать, что один двигатель перемещает щуп в направлении восток — запад и обратно, а другой в направлении север — юг и обратно. Задачей машины является проведение щупа по лабиринту к цели. Цель укреплена на шпильке, которая может быть вставлена по произволу в одно из гнезд, расположенных в центре каждого из двадцати пяти квадратов. Таким образом, условия задачи могут произвольно изменяться; сохраняется лишь размер лабиринта из 5×5 квадратов. После первого включения машина пытается найти решение лабиринтной задачи. Когда машина выключается, запоминающие реле по существу забывают все, что они запоминали, и после этого начинают работать заново, ничего не зная о лабиринте.

Севидж (Savage). Значит ли это, что реле находятся в нейтральном среднем положении, т. е. ни в правом, ни в левом?

¹⁾ Shannon C., Presentation of a maze-solving machine, Trans. of 8 Cybern. Conference, J. Macy, № 1, 1952, 173. [Стенограмма доклада.— Прим. перев.]



Р и с. 1

1—ролики каретки; 2—каретка; 3—гибкий кабель к каретке; 4—направляющие; 5—двигатель «восток-запад»; 6—двигатель «север-юг»; 7—центрирующий переключатель «восток-запад»; 8—управляющие переключатели «восток-запад»; 9—лампа; 10—чувствительный шуп; 11—перегородки лабиринта; 12—цель; 13—ходовой винт; 14—гнезда для цели; 15—индикаторные лампы; 16—счетчик шагов; 17—одношаговый режим; 18—автоматический режим; 19—ручной режим; 20—кнопки управления; 21—ручное управление двигателя; 22—выключатель.

Шеннон. Они находятся в некоторой исходной позиции, которая не является точно нейтральной позицией, но остается малозначащей. Теперь щуп исследует лабиринт в поисках цели. Когда он достигает центра квадрата, машина принимает новое решение испробовать следующее направление. Если щуп коснется перегородки, двигатель реверсируется, перемещая щуп вновь к центру квадрата, где выбирается новое направление движения. Выбор базируется на предыдущих сведениях и соответствует определенной «стратегии», которая не очень сложна.

Питтс (Pitts). Это стратегия фиксированная, а не случайная?

Шеннон. Да, здесь нет случайного элемента. Вначале предполагалось использовать вероятностный элемент, но потом было решено, что проще иметь дело с чистой стратегией. Когда щуп в своем исследовании достигает цели, двигатели останавливаются, на щупе зажигается лампочка и включается звонок. Машина решила лабиринтную задачу. Теперь можно вручную возвратить щуп обратно в начальное положение, и очевидно, что машина запомнила найденное ею решение, так как, когда машина снова включается, щуп пойдет прямо к цели, не натываясь на перегородки и минуя тупики. Он в состоянии пройти прямо к цели из любой части лабиринта, которую он посетил при исследовании. Если теперь передвинуть щуп в ту часть лабиринта, которую он еще не исследовал, он будет продолжать поиск и перемещаться до тех пор, пока не достигнет уже исследованной им зоны. Оттуда он пойдет прямо к цели.

Теперь хотелось бы указать на следующую особенность машины. Можно изменить лабиринт так, что решение, найденное машиной, перестанет действовать. Соответствующей перестановкой перегородок можно получить интересный эффект. В предшествующем лабиринте решение начинается в квадрате *A*, ведет в квадрат *B*, затем в *C* и далее к цели. Перестановкой перегородок можно вынудить машину из квадрата *C* отправиться в новый квадрат, квадрат *D*, и оттуда обратно к исходному квадрату *A*. Когда щуп приходит в квадрат *A*, он «вспоминает», что по старой программе он должен идти в квадрат *B*; в результате он начинает совершать путь последовательно по квадратам *A, B, C, D, A, B, C, D, ...* Устанавливается порочный круг или заикливание.

Джерард (Gerard). Невроз?

Шеннон. Да.

Севидж. Машина не может поступать так, если ее память ничего не хранит, но она может так поступать, если ее поведение обусловлено?

Шеннон. Да, только после того, как это поведение обусловлено заранее. Однако машина имеет антиневротические цепи, предназначенные для предотвращения подобной ситуации.

Мид (Mead). После того как она повторится несколько раз?

Шеннон. После того как машина повторит цикл около 6 раз, он будет прерван. Релейная схема, содержащая счетчик, прекращает цикл на двадцать четвертой попытке.

Франк (Frank). Сколько всего реле в машине?

Шеннон. В общей сложности около 75 реле.

Севидж. Не имеется ли какого-либо способа определить, что имеет место «психическое» расстройство, или только определяется, что процесс перемещения щупа совершается слишком долго?

Шеннон. Да. После выхода из цикла машина возвращается к исследовательской стратегии.

Тойбер (Teuber). Должна ли она снова изучить лабиринт полностью или она может использовать какую-либо информацию о его форме, полученную ранее?

Шеннон. Нет, если машина остановится, она начинает с начала и не может использовать полученную ранее информацию.

Севидж. Но я думаю, она пытается использовать ее. Она двигается, как будто она уже знает, куда следует двигаться.

Шеннон. На самом же деле вся старая информация начинает ей вредить.

Бигелу (Bigelou). Я думаю это так.

Шеннон. Да, она постепенно начинает продвигаться в направлении цели. Я хотел бы использовать оставшееся у меня время для того, чтобы пояснить некоторые особенности работы машины.

Стратегия, используемая машиной, может быть описана следующим образом: имеется два типа поведения машины, которые я называю «исследовательская стратегия» и «стратегия цели». Обе они достаточно просты. Исследовательская стратегия используется при первых попытках найти цель. Каждому квадрату лабиринта соответствует запоминающее устройство, состоящее из двух реле. Они могут запомнить одно из четырех возможных направлений: север, восток, юг или запад. При этом запоминается то направление, в котором щуп покинул данный квадрат в последнее посещение. В этом состоят все сведения, которые машина запоминает по пути щупа в лабиринте. Существуют некоторые другие функции памяти в вычислительной части схемы, но эти запомненные направления являются теми сведениями, которые позднее позволяют воспроизвести маршрут щупа.

Представим теперь запомненное направление в некотором квадрате D в виде вектора. В исследовательской стратегии, когда щуп входит в квадрат, машина берет вектор D и в качестве первого выбора поворачивает его на 90° . Пусть, например, щуп при своем последнем посещении покинул квадрат в восточном направлении. Если он снова попадет в этот же квадрат, он в качестве первой попытки будет испытывать северное направление. Если

щуп коснется перегородки и вернется обратно, он снова повернется на 90° , так как он только что ввел это северное направление в запоминающее устройство, и будет испытывать западное направление и т. д. Попытки следуют против часовой стрелки, начиная с того направления, в котором щуп покинул квадрат при своем последнем посещении, за одним исключением: щуп запоминает также направление, по которому он вошел в квадрат при данном посещении, но при первом повороте вектора D он пропускает это входное направление. Это предотвращает многократные повторения в траектории щупа. До того как эта особенность была установлена, имелась тенденция продолжать исследование до нового квадрата, возвращаться обратно через весь лабиринт, а затем выбирать следующий квадрат и т. д.; все это требовало очень много времени для решения лабиринтной задачи (приблизительно в три раза больше, чем сейчас).

Когда щуп достигнет цели, сработает реле, замыкающее цепи, и машина будет действовать далее по принципу «стратегии цели», которая также основывается на использовании вектора D .

В стратегии цели машина сразу использует направление D , в котором щуп покинул квадрат в свое последнее посещение. Это очень просто реализуется и очень удобно для решения лабиринтной задачи, так как исключает все тупики и кружения на одном месте. Так как тупик должен быть покинут щупом через тот же квадрат, через который он вошел, направление D , запомненное для данного квадрата, неизбежно поведет щуп к цели скорее, чем при заходе в тупик. Аналогично в случае, если машина следует по замкнутому (циклическому) пути или по незнакомой части, отыскивая путь к цели, направление, запомненное для последнего разветвления пути, должно вести к цели скорее, чем боковые ветви. Следовательно, машина (щуп) следует к цели по кратчайшему пути, который она уже установила.

Последняя особенность — способность забывать — заключается в следующем: предположим, что после достижения цели щуп переставлен на другое место лабиринта и поиск начался снова. Тогда машина начнет считать число сделанных движений щупа и если он не достигнет цели в течение определенного количества шагов, в данном случае двадцати четырех, машина устанавливает, что лабиринт был изменен или что щуп совершает движение по замкнутому циклу, или что-нибудь в этом роде, и следовательно, предыдущая программа уже не действует. Схема после этого переключается на исследовательскую стратегию, которая математически гарантирует нахождение пути в любом лабиринте, имеющем конечное решение.

Существуют также и другие стороны работы машины, представляющие некоторый интерес. Запоминающее устройство почти

не различает смысла запоминаемого; так, можно взять группу проводов, идущих от основания системы к запоминающему устройству, и поменять местами группы, соответствующие направлению «север — юг» и направлению «восток — запад», и машина при этом будет продолжать работать правильно и без заметных изменений, хотя данные, соответствующие некоторому квадрату, будут храниться в других частях запоминающего устройства.

Следующая особенность машины — это, конечно, наличие в ней большого количества цепей обратной связи. Наиболее значительной из них является цепь обратной связи от щупа через всю систему к двигателям и в виде механического перемещения обратно от двигателя к щупу. Обычно если имеется цепь обратной связи, то изменение знака обратной связи полностью нарушает действие всей системы. Существует большая разница между положительной и отрицательной обратной связью. Однако эта машина, решающая лабиринтную задачу, является такой, что можно менять как один, так и оба знака обратной связи, и она будет работать одинаково хорошо.

Это значит, что внутри данной системы понятия «правое» и «левое» взаимозаменяемы; другими словами, воздействие этого явления на стратегию машины заключается в том, что если одна из цепей обратной связи меняет знак, то это означает, что поворот на 90° против часовой стрелки заменяется таким же поворотом по часовой стрелке. Если знаки меняются в обеих цепях, стратегия не изменяется.

Фон Фёрстер (von Foerster). Если существуют два различных пути достижения цели, машина, конечно, способна выбрать только один. Имеется ли у нее при этом возможность выбрать наилучший путь?

Шеннон. Нет, она не обязательно выбирает лучший путь, хотя существует вероятность предпочтения ею более короткого при выборе из двух путей. Кроме того, исследовательская стратегия этой машины позволяет находить решение для любого лабиринта как односвязного, так и многосвязного. Некоторые классические программы решения лабиринтных задач пригодны лишь для односвязных лабиринтов. Примером этого может служить метод правой руки, когда при проходе по лабиринту все время держатся за его правую стенку. Этот способ позволяет решить любой простой лабиринт, но часто оказывается безуспешным, если в лабиринте попадаются замкнутые петли.

Севидж. Не появляется ли указанное явление цикличности вследствие того, что машина не находится в условиях действительного поиска?

Шеннон. Нет, это скорее наблюдается при применении стратегии цели, чем при исследовательской стратегии.

Севидж. Стратегия цели заключается в перемещении по тому пути, по которому щуп прошел ранее, но что должна делать машина, если попытка оказывается тщетной?

Шеннон. Щуп возвращается к центру квадрата, поворачивается на 90° и пробует это направление. Машина при этом остается в состоянии стратегии цели.

Севидж. Понимаю. Когда щуп попадет в следующий квадрат, он пробует идти дальше по уже знакомому направлению?

Шеннон. Правильно. Дело в том, что при поиске решения может быть изучена большая часть лабиринта, но не весь лабиринт полностью. Если же поместить щуп в квадрат, который он не посещал, он будет продолжать поиск методом испытаний и ошибок до тех пор, пока не достигнет знакомого квадрата, откуда и направится прямо к цели. Прежде неизвестные квадраты при этом добавляются к первоначальному решению.

Бигелу. Можно создать новую петлю на любом известном пути; щуп исследует эту новую петлю немедленно и не делает ошибки. Так ли это?

Шеннон. Правильно.

Бигелу. Потому что, когда щуп возвращается на основной путь, поиск идет в верном направлении, если он узнает соответствующий квадрат?

Шеннон. Я не убежден, что понял, что вы имеете в виду.

Бигелу. Щуп совершает направленный путь. Теперь, если предложить новый маршрут, который перенесет щуп с известного пути в неизвестную ему зону, он вновь вернется на правильный путь.

Шеннон. Такой боковой путь полностью исключается, так как действует стратегия цели.

Бигелу. Но если все начинается где-то внутри данной системы, то процесс идет правильно, начиная от места старта?

Шеннон. Это в случае стратегии цели, но не в исследовательской стратегии.

Бигелу. Что следует сделать для уменьшения времени поиска решения для того, чтобы обучение путем повторных попыток позволило найти наиболее короткий путь в сложном лабиринте?

Шеннон. Думаю, что это потребовало бы значительно большей памяти в виде реле в связи с необходимостью сохранения в ней разных решений лабиринтной задачи, а также дополнительных вычислительных релейных схем для сравнения и оценки этих решений. Это действительно следовало бы сделать, но это будет очень сложно; это означало бы создание машины гораздо более сложной, чем имеющаяся.

Севидж. К тому же нужно было бы решать, когда следует осуществлять поиски нового пути. Это действительно весьма

важная проблема, возникающая в любом случае фактического обучения человека. Если вы уже можете чистить картофель, почему вы должны заботиться о поиске лучшего способа чистки? Может быть, вы уже чистите правильно? Как вы это узнаете?

Фон Фёрстер. Что происходит, если в лабиринте не оказывается цели?

Шеннон. Если цели нет, машина периодически повторяет путь, разыскивая цель; это значит, что она прорабатывает путь, идущий через каждый квадрат, и исследует каждую перегородку, и если она не находит цели, маршрут повторяется снова и снова. Машина продолжает искать цель в каждом квадрате, чтобы быть уверенной, что она обследовала все квадраты.

Франк. Это все слишком по-человечески.

Брозин (Brosin). Джордж Оруэлл, покойный автор книги «1984»¹⁾, должен был бы это рассмотреть.

Фон Фёрстер. А после этого? Что произойдет, если установить цель на пути после того, как машина начнет такое периодическое движение?

Шеннон. Когда щуп коснется цели, машина остановится и перейдет к стратегии цели, после чего каждый очередной раз она будет выходить на эту цель. Кстати, рассмотрим вопрос с математической точки зрения: для каждого из 25 квадратов запоминающее устройство машины сохраняет в памяти векторные направления — север, восток, юг и запад. Таким образом, в целом запоминающее устройство содержит векторное поле, определенное на лабиринте из 25 клеток. Когда щуп движется по лабиринту, он непрерывно исправляет это запомненное векторное поле таким образом, что векторы направляются вдоль возможных путей в лабиринте, ведущих к месту нахождения щупа в каждый данный момент.

Тойбер. Если повернуть поле на 180° , будет ли оно продолжать правильно функционировать?

Мак-Каллок. Если реверсировать соединения и пустить двигатель в таком состоянии, что направление его вращения изменится, сможет ли при этом машина, как и прежде, найти путь к цели?

Шеннон. Только если переключить те же переключатели в машине, которые сообщают, какой квадрат занят (щупом) в данный момент. Если реверсировать двигатели, то надо для компенсации переключить и эти реле. Иначе щуп будет считать, что он движется одним путем и введет его в запоминающее устройство, а фактически он двигался в другом направлении.

¹⁾ Orwell G., «1984», Harcourt, Brace & Co, 1949; Signet Books, 798, 1950.

Д ж е р а р д. Это было бы похоже на перекрестное сшивание двигательных нервов у животного и появление вследствие этого сгибания вместо разгибания.

Б и г е л о у. Как вы думаете, трудно было бы создать систему, которая, вместо того чтобы забывать, возвращалась бы к началу и, помня все, что она делала в первом квадрате, попыталась бы сделать что-нибудь еще, например вести поиск в обратной последовательности? Если это не приведет к новому решению, вернуться назад и двигаться в обратной последовательности из второго квадрата, что потребует замены в запоминающем устройстве каждого хранящегося в памяти квадрата, через который проходит щуп. Другими словами, это потребует очень небольшого дополнения к запоминающему устройству, поскольку необходимо запомнить только одну последнюю картину (лабиринт). Но когда стратегия цели больше не приводит к решению (это определяется, когда число попыток превысило некоторое число N), можно использовать переключающие схемы, которые возвращают все к началу, и затем повторными проверками искать новое решение.

Ш е н н о н. Этот вопрос не рассматривался, но думаю, что все делалось бы слишком медленно, поскольку при этом большая часть времени уходила бы на возвратные движения, возвращения к началу и исследование разных гипотез.

Б и г е л о у. Если известно, как идти от начала к цели, всегда ли известно, как возвращаться от цели к началу с помощью простого реверса переключателей?

Ш е н н о н. Нет. Это векторное поле, если хотите, является уникальным при движении в направлении векторов, но при обратном движении появляются точки ветвления, так что машина не знает, откуда она пришла.

С е в и д ж. Ориентировано ли это поле в направлении к цели в каждой своей точке?

Ш е н н о н. Да. Если следовать по направлению векторов, цель будет достигнута, но в обратном направлении можно попасть в точку ветвления, из которой можно идти в любом направлении. В этом случае, изучая запоминающее устройство, нельзя сказать, откуда пришел щуп.

С е в и д ж. Эта неорганизованность, существующая в окрестности любой определенной начальной точки, является одной из особенностей машины. После того как она исследовала лабиринт, если начать работу в точке, где щуп уже был, она продолжает путь; если начать путь там, где щуп еще не был, он сначала найдет место, где он был, а затем будет продолжать путь к цели.

М а к К а л л о к. Подобно тому, как человек, знающий город, может из какого-нибудь места прийти в любое другое, но не всегда помнит, каким путем он шел.

ВКЛАД ФОН НЕЙМАНА В ТЕОРИЮ АВТОМАТОВ¹⁾

Теория автоматов возникла сравнительно недавно и, без сомнения, принадлежит к числу наиболее интенсивно развивающихся областей исследования. Она представляет собой науку, граничащую в математике с символической логикой и теорией машин Тьюринга, в инженерном деле — с теорией и применением универсальных вычислительных машин, в особенности к общим проблемам неарифметического характера, а в биологии — с нейрофизиологией, теорией нервных сетей и т. д. Теория автоматов охватывает различные проблемы, начиная с проблем «геделевского типа» (относящихся к машинам Тьюринга и проблемам разрешения) и кончая проблемами размножения, приспособления, самовоспроизведения и самовосстановления и др. применительно к машинам.

Последние несколько лет своей жизни фон Нейман много работал в области теории автоматов, которая соединила его ранние исследования в логике и теории доказательств и его более поздние работы времени второй мировой войны и послевоенного периода, относящиеся к области универсальных вычислительных машин.

Теория автоматов, включающая в себя элементы «чистой» и прикладной математики наряду с элементами других наук, была идеальным полем для разностороннего интеллекта фон Неймана. Он внес в эту теорию много новых идей и положил начало по меньшей мере двум новым направлениям исследования. К сожалению, он не смог довести до конца начатую работу, часть которой осталась в черновиках и неизданных лекциях, а часть нигде не записана и восстановить ее можно только по воспоминаниям его коллег о случайных разговорах.

Здесь не будет рассматриваться его огромной важности вклад в теорию вычислительных машин и их применение; его идеи, касаю-

¹⁾ Shannon C., Von Neumann's contributions to automata theory, *Bull. Amer. Math. Soc.*, 64, № 2 (1958), 123.

щиеся логической структуры машин, использования блок-схем для программирования, и методов программирования различных задач, таких, как обращение матриц, метод Монте-Карло¹⁾ и т. д.—ограничимся областью собственно теории автоматов.

Надежные машины и ненадежные элементы. Одна из важных частей работы, проделанной фон Нейманом в теории автоматов, относится к проблеме синтеза надежных машин из ненадежных элементов²⁾.

Пусть дано множество элементарных блоков с некоторыми положительными вероятностями неправильного функционирования. Можно ли из этих блоков при помощи соответствующего метода синтеза строить произвольно большие и сложные автоматы, для которых вероятность появления ошибки на выходе поддавалась бы контролю? Можно ли сделать вероятность ошибки сколь угодно малой или хотя бы не превосходящей некоторого фиксированного значения (не зависящего от конкретного автомата)?

Мозг человека и животных дает нам пример очень большой и относительно надежной системы, построенной из индивидуальных компонент, нейронов, которые ненадежны не только в выполнении операций, но и в тонких деталях взаимосвязи. Более того, хорошо известно, что при повреждении, несчастном случае, болезни и т. д. мозг продолжает функционировать замечательно правильно, даже если поражены его большие области.

Эти факты представляют сильный контраст по сравнению с поведением и организацией современных вычислительных машин. Индивидуальные элементы этих машин должны быть выполнены с чрезвычайной надежностью, каждый провод должен быть соединен нужным образом и каждая команда в программе должна быть правильной. Любая ошибка в элементе, в соединении элементов или в программе обычно приводит к полному искажению результатов. Если рассматривать мозг как машину, то очевидно, что предохранение от ошибок организовано в нем совершенно иначе, чем в вычислительных машинах.

¹⁾ von Neumann J., Burks A., Goldstine H., Preliminary discussion of the logical design of an electronic computing instrument, Report prepared for U. S. Army Ord. Dept. under contract W-36-034-ORD-7481, part I, June 28 (1942); 2d ed. Sept. 2 (1947), von Neumann J., Goldstine H., Numerical inverting of matrices of high order, *Amer. Math. Soc. Bull.*, 53 (1947), 1021. von Neumann J., Goldstine H., Planning and coding of problems for an electronic computing instrument, Report prepared for U. S. Army Ord. Dept. under contract W-36-034-ORD-7481, I, II and III, part II (1947), 69, 68 and 23.

²⁾ von Neumann J., Probabilistic logics and the synthesis of reliable organisms from unreliable components, «Automata studies», edited by Shannon C., MacCarthy J., Princeton University Press, 1956, 43; русский перевод в сб. Автоматы, ИЛ, М., 1956.

Эта проблема аналогична проблеме, возникающей в теории связи, когда требуется построить такие коды для передачи информации, что надежность полного кода высока даже в тех случаях, когда надежность передачи отдельных символов мала. В теории связи эту проблему можно решить соответствующим введением избыточности и в данном случае нужно применить аналогичные приемы. Здесь недостаточно простого выполнения одних и тех же вычислений много раз подряд и выбора значения по большинству. Значение по большинству берется от ненадежных элементов, и так много раз подряд — значение по большинству от значений по большинству и т. д. Здесь возникает ситуация: «кто будет сторожить сторожа».

Исследование этих проблем фон Нейман начал с рассмотрения формальной структуры автомата. Та система, которую он выбрал, аналогична модели Мак Каллока—Питтса; схемы состоят из отдельных элементов относительно простого типа, связанных между собой. Каждый элемент получает двоичные сигналы на входы от множества различных входных линий и выдает выходные двоичные сигналы на некоторую выходную линию. Сигнал на выходе появляется через целое число единиц времени после подачи сигнала на вход. Если бы выходной сигнал был функцией значений входных сигналов, имелся бы надежный элемент, который может выполнять операцию «и», «не», штрих Шеффера и т. д. Однако если выходной сигнал зависит от входных только статистически, например с вероятностью $1-\epsilon$, на выходе получается штрих Шеффера и с вероятностью ϵ — отрицание этой операции, то имеется ненадежный элемент. Если же дано неограниченное число таких ненадежных элементов, например элементов для реализации штриха Шеффера, то можно ли из них построить надежный вариант любого заданного автомата?

Фон Нейман показал, что это можно сделать, и проиллюстрировал это двумя совершенно различными приемами. Первый из них, возможно, более красив математически, так как он тесно связан с описанной проблемой и близко подходит к проблеме «сторожа».

Решение состоит в конструировании из трех ненадежных под-схем и некоторых сравнивающих устройств одной более крупной и более надежной подсхемы, выполняющей ту же функцию, что и исходная подсхема. Прodelывая это для каждого элемента схемы с ненадежными элементами, получим схему с тем же поведением, что и у заданной, но состоящую из ненадежных элементов.

Первый прием, как он указывал, страдает двумя недостатками. Во-первых, окончательная надежность не может быть сделана произвольно высокой, а может быть доведена только до определенного уровня ϵ (ϵ зависит от надежности исходных элементов). Если эти элементы очень низкого качества, то решение едва ли может

считаться удовлетворительным. Во-вторых, что более важно с практической точки зрения, требуемая избыточность в большинстве случаев фантастически велика. Число требуемых элементов растет экспоненциально по отношению к числу n элементов, необходимых для создания моделируемого автомата. Так как во всех случаях, представляющих практический интерес, n очень велико, то это решение ценно только с точки зрения логической возможности.

Второй прием состоит в том, что фон Нейман называл «мульти-триоком». Он заключается в том, что двоичный выход в машине представляется не одной линией, а пучком из N линий, и двоичный выходной сигнал определяется в зависимости от того, много линий или, наоборот, очень мало линий несут значения 1. Метод синтеза автоматов, основанный на использовании надежных элементов, в этом случае заменяется методом, в котором каждая линия становится пучком линий, а каждый элемент заменяется подсхемой, которая оперирует соответствующим образом с пучками входных и выходных линий. Фон Нейман показал, каким образом можно сконструировать такие подсхемы. Он также сделал некоторые оценки избыточности, требуемой для достижения определенной надежности. Например, вместо одного ненадежного «мажоритарного» элемента, вероятность ошибки которого равна $1/200$, использованием избыточности в $60\,000$ к 1 можно построить подсхему, представляющую мажоритарный элемент для пучков и с вероятностью ошибки 10^{-20} . Произведя соответствующий подсчет, увидим, что этот автомат, обладающий сложностью и быстродействием мозга, может работать в течение ста лет, сделав при этом всего несколько ошибок. Другими словами, нечто родственное этой схеме может обладать по крайней мере такой же надежностью, как мозг.

Самовоспроизводящиеся машины. Другой ветвью теории автоматов, которую развивал фон Нейман, является изучение проблемы самовоспроизведения машин или проблемы — можно ли построить простую и абстрактную систему «машин», которые способны строить другие идентичные машины или даже способные к некоторого рода эволюционному процессу, в котором последующие поколения строят машины более высокой «сложности». Реальная трудность здесь состоит в том, чтобы соответствующим образом сбалансировать простоту формальных построений и легкость обращения, с одной стороны, и степень близости модели к реальным физическим машинам, с другой стороны. Если модель слишком близка к реальности, нам приходится кодировать все сложнейшие аспекты природы, большая часть которых не имеет никакого отношения к вопросу самовоспроизведения. Однако при слишком сильном упрощении модель становится настолько абстрактной и упрощенной, что проблема выглядит почти тривиальной и решение не дает ничего нового с точки зрения выяснения поставленного философского вопроса.

Однажды после продолжительной дискуссии о трудности удовлетворительно сформулировать эту проблему фон Нейман заметил: «Я не собираюсь всерьез возражать тем, кто утверждает: а) каждому известно, что автомат может воспроизвести сам себя, и б) каждому известно, что он этого сделать не может».

Фон Нейман посвятил много времени исследованию вопроса о самовоспроизведении автоматов; краткое изложение его высказываний по этому поводу содержится в материалах Хиксонского симпозиума¹⁾, и подробное — в более поздних незаконченных рукописях²⁾.

Фактически он рассмотрел две различные формулировки проблемы. Модель, о которой он говорил в докладе на Хиксонском симпозиуме и в более ранних работах на эту же тему, состоит из небольшого числа различных исходных элементов. Этими элементами могут быть, например, стойка, датчик (для установления присутствия других частей), соединительный элемент (для скрепления отдельных частей вместе) и т. д. Машины представляют собой комбинации этих частей и существуют в геометрической среде вместе с другими подобными частями, которые могут быть взяты машиной из этой окружающей среды.

Некоторые машины, собранные из таких частей, способны группировать вместе и оперировать отдельными частями из окружающей среды. Можно также сконструировать машину с «программным» управлением, которая выполняет такие же длинные последовательности команд, как и обычная вычислительная машина. Здесь, однако, команды содержат операции над такими отдельными частями, а не над числами в ходе длинных вычислений. Ситуация несколько аналогична той, которая имеет место в случае машин Тьюринга, и действительно существует понятие *универсальной конструирующей машины*, которая может при наличии соответствующей программы моделировать любую машину. Фон Нейман показал, что такая универсальная машина вместе с устройством, копирующим программы, может быть превращена в самовоспроизводящуюся машину.

Эта модель очень интересна сама по себе, но так как она требует изучения сложных случаев движения частей в действительном евкли-

¹⁾ von Neumann J., The general and logical theory of automata, Cerebral mechanisms in behavior — The Hixon symposium September 1948, Pasadena, ed. by Jeffres L., John Wiley and Sons, Inc., New York, 1951, 1—31. Русский перевод: фон Нейман Дж., Общая и логическая теория автоматов. Приложение к книге Тьюринга А., Может ли машина мыслить, Физматгиз, М., 1960.

²⁾ von Neumann J., Незаконченная рукопись работы для издательства Иллинойского университета по теории автоматов и, в частности, по самовоспроизводящимся машинам.

довом пространстве, то выполнить ее во всех подробностях чрезвычайно трудно, даже если не учитывать проблем источников питания, шумов в среде и тому подобное. Во всяком случае, фон Нейман в своих более поздних работах перешел от этой модели к более простой.

Второй тип самовоспроизводящейся системы описан в неоконченной книге, предназначенной для издательства Иллинойского университета. Эта вторая модель, возможно, больше напоминает биологическое размножение на низшем уровне (скажем, на клеточном или даже молекулярном), хотя она и не следует в точности никакой реальной физической системе.

Рассмотрим бесконечное множество квадратов на евклидовой плоскости, каждый квадрат, или «клетка», может находиться в некотором числе состояний. В модели, построенной фон Нейманом, клетки имеют двадцать девять состояний. Время пробегает дискретные моменты. Состояние клетки в данный момент есть функция ее состояния в предыдущий момент и состояний четырех ее ближайших соседей в предыдущий момент. С течением времени состояния всех клеток изменяются в соответствии с этими функциональными соотношениями. Определенное состояние клеток называется «начальным» и соответствует неактивной части плоскости. При соответствующем подборе функциональных уравнений можно получить группы соседних клеток, которые ведут себя подобно живому организму, способному сохранять свою индивидуальность, передвигаться и даже воспроизводить себя в смысле порождения других групп клеток с теми же активными состояниями.

В дополнение к вопросу о самовоспроизведении автоматов фон Нейман в некоторой степени касался проблем «эволюции» в автоматах или проблемы — можно ли построить автоматы, которые будут конструировать последовательные поколения автоматов, могущих в некотором смысле приспосабливаться к окружающей среде. Он указывал на существование критических размеров автомата, построенного из элементов заданного типа, а именно что автоматы, размер которых меньше критического, могут конструировать только автоматы меньшие, чем они сами, в то время, как некоторые автоматы, размер которых равен или больше критического, способны к самовоспроизведению и даже к эволюции (при соответствующем задании условий и требований приспособляемости).

Сравнение вычислительных машин и мозга. Огромный интерес для фон Неймана представлял вопрос о соотношении между центральной нервной системой и современными универсальными вычислительными машинами. В его докладе на Хиксонском симпозиуме этой теме уделено не меньше внимания, чем проблеме самовоспроизводящихся машин. Более широко эти вопросы затронуты

в лекциях на Силлимановских чтениях (которые он подготовил, но не смог прочитать).¹⁾

Выявляя аналогии между вычислительными машинами и нервными сетями, фон Нейман ясно понимал и часто подчеркивал существенные различия между ними. При поверхностном рассмотрении очевидны различия в порядке числа и размеров их элементов и в скорости выполнения ими операций. Нейроны мозга работают гораздо медленнее, чем их искусственные аналоги — транзисторы или электронные лампы, но, с другой стороны, они имеют намного меньшие размеры, расходуют меньше энергии и число их в мозгу на несколько порядков больше, чем в самых больших вычислительных машинах. При более глубоком изучении обнаруживаются, как подчеркивал фон Нейман, различия в логической организации, определяющей функционирование этих двух типов систем. Отчасти эти различия определяются характером решаемых задач, который проявляется в «логической глубине», т. е. в числе элементарных операций, которые должны быть последовательно выполнены для получения решения. В вычислительных машинах это число может иметь порядок 10^7 или больше из-за определенной искусственности и «последовательности» метода решения задач. Мозг, обладая большим числом более медленных элементов, по-видимому, работает ближе к параллельному методу при меньшей логической глубине, и, кроме того, задачи, которые он решает, требуют меньшего разнообразия последовательных вычислений.

В лекциях на Силлимановских чтениях фон Нейман кратко коснулся любопытной идеи, относящейся к основаниям математики. Тьюринг в своей хорошо известной работе о вычислимости показал, как одна вычислительная машина может моделировать другую. Команды для второй машины переводятся с помощью сокращенной программы в последовательности команд для первой машины, что заставляет ее выполнять (в общем случае косвенным образом) то, что должна делать вторая машина. С помощью таких переводных программ можно заставить первую машину в конечном итоге делать то же, что вторая машина, хотя в действительности она работает в ином внутреннем коде. Этот прием стал обычным и весьма полезным в повседневном использовании вычислительных машин.

Если принять мозг как некую разновидность вычислительной машины, то вполне возможно предположить, что внешний язык, которым мы пользуемся при общении друг с другом, может быть совершенно отличен от внутреннего языка, используемого при вычислениях (которые включают, разумеется, все логические

¹⁾ von Neumann J., The Silliman memorial lectures, Yale University Press, 1958. Русский перевод: фон Нейман Дж., Вычислительная машина и мозг, Кибернетический сб., вып. 1, ИЛ, М., 1960, стр. 11.

и информационные операции наряду с арифметическими вычислениями).

Фон Нейман действительно приводит убедительные аргументы, говорящие о том, что до сих пор нет совершенно никакого представления о природе того первичного языка, который используется при работе мозга. Он говорит: «Логика и математика в центральной нервной системе, рассматриваемые как некоторые языки, должны иметь структуру, существенно отличную от структуры языков, к которым приводит нас повседневный опыт.

Нужно также отметить, что тот язык, о котором идет речь, может скорее соответствовать сокращенным программам в том смысле, как это было сказано выше, чем полным программам: когда мы говорим о математике, то при этом имеется в виду некоторый *вторичный* язык, построенный на основе *первичного* языка, используемого в действительности центральной нервной системой. Таким образом, внешняя форма нашей математики не является абсолютно существенной с точки зрения выяснения вопроса, что представляет собой математический или логический язык, в действительности используемый центральной нервной системой. Однако сделанные выше замечания относительно надежности, логической и арифметической глубины показывают, что эта система, чем бы она ни была, не может не отличаться существенным образом от того, что мы сознательно и явно рассматриваем как математику».

В целом вклад фон Неймана в теорию автоматов, так же как и его вклад в другие разделы математики и другие науки, характеризуется открытием совершенно новых областей исследования и прозорливостью в оценке возможностей приложений методов современной математики. Исследования в направлениях, открытые им для разработки, не будут полностью закончены еще в течение долгих лет. Очень жаль, что часть его замыслов относительно теории автоматов осталась нереализованной.

Теория информации

Введение

Современное развитие различных методов модуляции, таких, как ИКМ и ИФМ²⁾, позволяющих компенсировать изменение отношения сигнал/шум изменением полосы частот, повысило интерес к общей теории связи. Некоторые основные положения этой теории содержатся в важных работах Найквиста³⁾ и Хартли⁴⁾. В настоящей статье мы расширим теорию с тем, чтобы включить некоторое число новых факторов, в частности, влияние шума в канале и возможность экономии за счет учета статистической структуры исходного сообщения и назначения передаваемой информации.

Основная задача связи состоит в точном или приближенном воспроизведении в некотором месте сообщения, выбранного для передачи в другом месте. Часто сообщения имеют значение, т. е. относятся к некоторой системе, имеющей определенную физическую или умозрительную сущность, или находятся в соответствии с некоторой системой. Эти семантические аспекты связи не имеют отношения к технической стороне вопроса. Существенно лишь, что посылаемое сообщение является сообщением, *выбранным из некоторого*

¹⁾ A mathematical theory of communication, *Bell System Techn. J.*, 27 (1948), № 3, 379—423; 27 (1948), № 4, 623—656.

²⁾ ИКМ и ИФМ — сокращенные обозначения для импульсно-кодовой и импульсно-фазовой модуляции. Подробнее см., например, Харкевич А. А., Теоретические основы радиосвязи, М., 1947, гл. 2.— *Прим. ред.*

³⁾ Nyquist H., Certain factors affecting telegraph speed, *Bell Syst. Techn. J.*, April (1924), 324; Certain topics in telegraph transmission theory, *AIEE Trans.*, 47 (1928), April 617.

⁴⁾ Hartley R. V. L., Transmission of information, *Bell Syst. Techn. J.* (1928), July, 535 [русский перевод: Хартли Р., Передача информации, в сборнике «Теория информации и ее приложения», Физматгиз, 1959, стр. 5—35. Эта работа была опубликована также в книге: Shannon C. E., Weaver E., Mathematical theory of communication, Univ. of Ill. Press, Urbana, Ill., 1949, в которую, кроме этой работы, входила популяризирующая ее статья Уивера. В связи с этим эту работу часто неправильно цитируют как совместную работу Шеннона и Уивера. — *Прим. ред.*]

множества возможных сообщений. Система связи должна быть спроектирована так, чтобы ее можно было использовать для передачи любого возможного сообщения, а не только того, которое будет в действительности выбрано, так как результат этого выбора еще неизвестен в момент проектирования.

Если множество возможных сообщений конечно, то число сообщений или любую монотонную функцию от этого числа можно рассматривать как меру информации, создаваемой выбором одного сообщения из этого множества, в предположении, что все сообщения равновероятны. Как было указано Хартли, наиболее естественно выбрать логарифмическую функцию. Хотя это определение должно быть значительно обобщено при рассмотрении влияния статистической структуры сообщения и при наличии непрерывного множества сообщений, будем по существу во всех случаях пользоваться логарифмической мерой.

Логарифмическая мера более удобна по различным причинам.

1) Она практически более пригодна. Параметры, имеющие техническое значение, такие, как время, ширина полосы частот, количество реле и т. д., зависят линейно от логарифма числа возможностей. Например, добавление одного реле к некоторой схеме удваивает число возможных состояний реле. Тем самым прибавляется единица к логарифму этого числа при основании 2. Удвоение времени, грубо говоря, возводит в квадрат число возможных сообщений, т. е. удваивает логарифм и т. д.

2) Она ближе к нашему интуитивному представлению о подходящей мере. Это обстоятельство тесно связано с первой причиной, так как мы интуитивно измеряем количества с помощью линейного сравнения с принятыми эталонами. Например, каждый чувствует, что две перфокарты должны обладать вдвое большей емкостью для хранения информации, чем одна, а два идентичных канала должны иметь удвоенную пропускную способность.

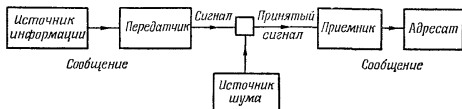
3) Она является более подходящей с математической точки зрения. Многие предельные переходы весьма просты при использовании логарифмов, но потребовали бы сложных приемов при использовании числа сообщений.

Выбор основания логарифмов соответствует выбору единицы измерения информации. Единицы измерения, получающиеся при использовании основания два, могут быть названы двоичными единицами или сокращенно битами (слово, предложенное Тьюки). Прибор с двумя устойчивыми состояниями, такой, как реле или мультивибратор, может хранить один бит информации. N таких устройств могут хранить N бит, так как полное число возможных состояний есть 2^N и $\log_2 2^N = N$. При использовании в качестве основания числа десять единицы измерения информации могут быть

названы десятичными единицами. Так как

$$\log_2 M = \log_{10} M / \log_{10} 2 = 3,32 \log_{10} M,$$

то десятичная единица составляет примерно $3 \frac{1}{3}$ бита. Числовой барабан арифмометра имеет десять устойчивых положений и поэтому обладает способностью хранить одну десятичную единицу информации. При аналитических расчетах, когда приходится интегрировать и дифференцировать, иногда удобно применять основание e .



Р и с. 1. Общая схема системы связи.

Получающиеся при этом единицы информации могут быть названы натуральными единицами. Переход от основания a к основанию b требует лишь умножения на $\log_b a$.

Под системой связи мы будем понимать систему типа, указанного на рис. 1. Она состоит по существу из пяти частей:

1. *Источник информации*, создающий сообщение или последовательность сообщений, которые должны быть переданы на приемный конец. Сообщения могут быть различных типов: а) последовательность букв, как в системах телеграфа и телетайпа; б) некоторая функция времени $f(t)$, как в радио или телефонии; в) функция времени и других переменных, как в черно-белом телевидении — здесь сообщение можно считать функцией $f(x, y, t)$ от двух пространственных координат и времени (интенсивность света в точке (x, y) и в момент времени t на экране приемной трубки); г) две или более функций времени, скажем $f(t), g(t), h(t)$, как это имеет место в «трехмерной» передаче звука или в системе, которая рассчитана на обслуживание нескольких индивидуальных каналов мультиплексным методом; д) несколько функций от нескольких переменных — в цветном телевидении сообщение состоит из трех функций: $f(x, y, t), g(x, y, t), h(x, y, t)$, определенных в трехмерном континууме. Мы можем также рассматривать эти три функции как компоненты векторного поля, определенного в этой области; аналогично несколько источников черно-белого телевидения производили бы «сообщения», состоящие из нескольких функций трех переменных; е) встречаются также различные комбинации, например, в телевидении с каналом звукового сопровождения.

2. *Передачик*, который перерабатывает некоторым образом сообщение в сигналы, соответствующие характеристикам данного канала. В телефонии эта операция состоит просто в преобразовании звукового сигнала в пропорциональный ему электрический ток. В телеграфии имеется некоторая операция кодирования, которая дает последовательность точек, тире и пробелов, соответствующих сообщению. В многоканальных системах ИКМ для создания сигнала нужно выбрать некоторые значения речевых функций, затем подвергнуть их сжатию, квантованию, кодированию и, наконец, соответствующему перемешиванию. Системы вокодера, телевидение и частотная модуляция являются другими примерами сложных операций, применяемых к сообщению для получения сигналов.

3. *Канал* — это среда, используемая для передачи сигнала от передатчика к приемнику. Каналом может быть пара проводов, коаксиальный кабель, полоса радиочастот, луч света и т. д.¹⁾

4. *Приемник* обычно выполняет операцию, обратную по отношению к операции, производимой передатчиком, восстанавливая сообщение по сигналам.

5. *Адресат* — это лицо (или аппарат), для которого предназначено сообщение.

Рассмотрим некоторые общие проблемы, относящиеся к системам связи. Для этого необходимо прежде всего представить различные элементы, входящие в эти системы, в виде математических понятий, идеализированных подходящим образом по сравнению с их физическими прообразами. Мы можем грубо разделить системы связи на три главные категории: дискретные, непрерывные и смешанные. Под дискретной системой будем понимать систему, в которой и сообщение и сигнал являются последовательностями дискретных символов. Типичным примером служит телеграфия, где сообщение представляется в виде последовательности букв, а сигнал — последовательности точек, тире и промежутков. Непрерывная система — это такая система, в которой и сообщение и сигнал рассматриваются как непрерывные функции, например радио или телевидение. Смешанная система — это такая система, в которой встречаются и дискретные и непрерывные переменные, например передача речи с помощью ИКМ. Рассмотрим сначала дискретную систему, — случай, имеющий применение не только в теории связи, но также в теории вычислительных машин, планировании телефонных станций и других областях. Кроме того, исследование дискретной системы образует основу для исследования непрерывной и смешанной систем, которые будут рассмотрены во второй части статьи.

¹⁾ В процессе передачи сигнал может быть искажен шумом. Это схематически показано на рис. 1 с помощью источника шума, действующего на передаваемый сигнал, в результате чего получается принимаемый сигнал.

I. ДИСКРЕТНЫЕ СИСТЕМЫ БЕЗ ШУМОВ

1. Дискретный канал без шума

Телетайп и телеграф являются двумя простыми примерами использования дискретного канала для передачи информации. Вообще же под дискретным каналом будет подразумеваться система, с помощью которой можно передать из одного места в другое последовательность символов, выбранных из некоторого конечного множества элементарных символов $S_1, \dots, S_n$. Предполагается, что каждый из символов S_i имеет определенную длительность во времени t_i секунд (не обязательно ту же самую для разных S_i , например точки и тире в телеграфии). Не требуется, чтобы все возможные последовательности символов S_i могли передаваться системой. Эти допустимые последовательности и будут возможными сигналами для канала. Так, например, в телеграфии такими сигналами являются: 1) точка, создаваемая замыканием линии на некоторое единичное время и последующим размыканием на такое же время; 2) тире, создаваемое замыканием на три единицы времени и размыканием на одну единицу; 3) пробел между буквами, создаваемый, скажем, размыканием на три единицы времени; 4) пробел между словами, создаваемый размыканием на шесть единиц времени. Мы могли бы наложить ограничение на допустимые последовательности, состоящее в том, чтобы пробелы не следовали один за другим (так как два последовательных пробела между буквами идентичны пробелу между словами). Рассмотрим теперь вопрос о том, как можно измерить пропускную способность такого канала.

В случае телетайпа, где все символы имеют одинаковую длительность и каждая последовательность из этих 32 символов допустима, ответ получить легко. Каждый символ несет пять битов информации. Если система передает n символов в секунду, естественно сказать, что канал имеет пропускную способность $5n$ битов в секунду. Это еще не означает, что канал телетайпа будет всегда передавать информацию с такой скоростью — это максимальная возможная скорость, и достигает ли или нет действительная скорость этого максимума, зависит, как будет показано ниже, от источника информации на входе этого канала.

Для более общего случая, когда длины символов различны и имеются ограничения на допустимые последовательности, дадим следующее определение: пропускная способность C дискретного канала задается формулой

$$C = \lim_{T \rightarrow \infty} \frac{\log N(T)}{T},$$

где $N(T)$ — число допустимых сигналов длительности T

Легко видеть, что в случае телетайпа это определение сводится к предыдущему. Можно показать, что рассматриваемый предел существует и для большинства случаев, представляющих интерес, конечен. Предположим, что все последовательности символов $S_1, \dots, S_n$ допустимы и что эти символы имеют длительности $t_1, \dots, t_n$. Чему равна пропускная способность канала? Если $N(t)$ представляет собой число последовательностей длительности t , то

$$N(t) = N(t - t_1) + N(t - t_2) + \dots + N(t - t_n).$$

Полное число всех последовательностей равно сумме чисел последовательностей, оканчивающихся на $S_1, \dots, S_n$; а эти числа равны $N(t - t_1), N(t - t_2), \dots, N(t - t_n)$, соответственно. Согласно хорошо известному результату исчисления конечных разностей, $N(t)$ будет при больших t асимптотически приближаться к AX_0^t , где A — константа и X_0 — наибольший действительный корень характеристического уравнения

$$X^{-t_1} + X^{-t_2} + \dots + X^{-t_n} = 1,$$

и поэтому

$$C = \lim_{T \rightarrow \infty} \frac{\log AX_0^T}{T} = \log X_0. \quad 1)$$

Если имеются ограничения на допустимые последовательности, часто можно все же получить разностное уравнение такого же типа и найти C из характеристического уравнения. В упомянутом выше случае телеграфии

$$N(t) = N(t - 2) + N(t - 4) + N(t - 5) + N(t - 7) + \\ + N(t - 8) + N(t - 10),$$

в чем легко убедиться, подсчитывая последовательности символов с учетом последнего или следующего за последним символом. Отсюда C равно $-\log \mu_0$, где μ_0 — положительный корень уравнения

$$1 = \mu^2 + \mu^4 + \mu^5 + \mu^7 + \mu^8 + \mu^{10}.$$

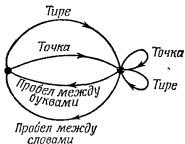
Решая это уравнение, находим, что $C = 0,539$.

Весьма общий тип ограничений, которые могут быть наложены на допустимые последовательности, состоит в следующем: предположим, что в системе имеется некоторое число возможных состояний $a_1, a_2, \dots, a_m$. В каждом состоянии могут быть переданы только некоторые символы из множества $S_1, \dots, S_n$ (различные подмножества для различных состояний). После того как один из символов передан, состояние меняется на некоторое новое состояние, зависящее как от старого состояния, так и от конкретного переданного символа. Простым примером этого служит телеграф. Здесь имеются

1) См., например, Гельфонд А. О., Исчисление конечных разностей, Гостехиздат, М., 1952, гл. V, § 4. — Прим. ред.

два состояния, зависящие от того, был ли последним переданным символом пробел или нет. Если это был пробел, то после него могут быть посланы только точка или тире и состояние всегда изменяется. Если нет, то может быть передан любой символ, и состояние изменяется, если посылается пробел, в противном случае оно остается неизменным. Эти условия могут быть представлены в виде схемы, как показано на рис. 2. Узловые точки соответствуют состояниям, а линии указывают символы, возможные в некотором состоянии, и результирующие состояния. В приложении 1 показано, что если ограничения на допустимые последовательности могут быть описаны в такой форме, то S существует и может быть вычислено согласно следующей теореме.

Теорема 1. Пусть $b_{ij}^{(s)}$ означает длительность s -го символа, который допустим в состоянии i и приводит к состоянию j .



Р и с. 2. Графическое представление телеграфных символов.

Тогда пропускная способность канала S равна $\log W$, где W — наибольший действительный корень характеристического уравнения

$$\left| \sum_s W^{-b_{ij}^{(s)}} - \delta_{ij} \right| = 0;$$

здесь $\delta_{ij} = 1$, если $i = j$, и нуль в противном случае.

Например, в случае телеграфа (рис. 2) уравнение имеет вид

$$\begin{vmatrix} -1 & W^{-2} + W^{-4} \\ W^{-3} + W^{-6} & W^{-2} + W^{-4} - 1 \end{vmatrix} = 0.$$

После разложения это уравнение сводится к уравнению, приведенному выше для рассматриваемого множества ограничений.

2. Дискретный источник сообщений

Как мы уже видели, при весьма общих условиях логарифм числа возможных сигналов в дискретном канале линейно возрастает со временем. Пропускная способность может быть охарак-

теризована заданием скорости этого возрастания, т. е. числа битов в секунду, требуемых для воспроизведения конкретного сигнала.

Рассмотрим теперь источник сообщений. Как следует математически описывать источник сообщений и какое количество информации, измеряемое числом битов в секунду, создается данным источником? Главная сторона этого вопроса заключается в использовании статистических сведений об источнике для уменьшения пропускной способности канала путем применения надлежащего метода кодирования информации. В телеграфии, например, сообщения, которые должны быть переданы, состоят из последовательностей букв. Эти последовательности, однако, не являются вполне случайными. Вообще говоря, они образуют фразы и имеют статистическую структуру, скажем, английского языка. Буква Е встречается чаще, чем Q; последовательность ТН чаще, чем ХР, и т. д.

Наличие такой структуры позволяет получить некоторую экономию времени (или пропускной способности канала) с помощью подходящего кодирования последовательностей сообщений в последовательности сигналов. В ограниченных пределах это уже делается в телеграфии путем использования наиболее короткого символа в канале — точки — для наиболее распространенной в английском языке буквы Е, в то время как редко встречающиеся буквы Q, X, Z представляются более длинными последовательностями точек и тире. Этот принцип проводится еще дальше в некоторых коммерческих кодах, где наиболее употребительные слова и фразы представляются четырех- или пятибуквенными кодовыми группами, что дает значительную экономию среднего времени. В стандартизированных и поздравительных телеграммах, используемых в настоящее время, этот принцип развивается еще дальше. Там проводится кодирование одного или двух предложений в относительно короткую последовательность цифр.

Можно считать, что дискретный источник создает сообщение символ за символом. Он будет выбирать последовательные символы в соответствии с некоторыми вероятностями, зависящими, вообще говоря, как от предыдущих выборов, так и от конкретного рассматриваемого символа. Физическая система или математическая модель системы, которая создает такую последовательность символов, определяемую некоторой заданной совокупностью вероятностей, называется вероятностным процессом¹⁾. Поэтому можно счи-

¹⁾ См., например, Chandrasekhar S., Stochastic problems in physics and astronomy, *Rev. of Modern Physics*, 15, № 1, January (1943), 1. Русский перевод: Чандрасекар С., Стохастические проблемы в физике и астрономии, ИЛ, М, 1947.

тять, что дискретный источник представляется некоторым вероятностным процессом. Обратно, любой вероятностный процесс, который создает дискретную последовательность символов, выбираемых из некоторого конечного множества, может рассматриваться как дискретный источник. Это включает такие случаи, как:

1. Печатные тексты на таких языках, как английский, немецкий, китайский.

2. Непрерывные источники сообщений, которые превращены в дискретные с помощью некоторого процесса квантования. Например, квантованная речь из ИКМ-передатчика или квантованный телевизионный сигнал.

3. Математические случаи, когда просто определяется абстрактно некоторый вероятностный процесс, который порождает последовательность символов. Приведем следующие примеры источников этого последнего типа:

А) пусть имеются пять букв А, В, С, D, Е, каждая из которых выбирается с вероятностью 0,2 независимо от результатов предыдущих выборов. Это привело бы к последовательностям, типичным примером которых является следующая последовательность:

BDCBCECCADCBDDAAECSEEAABBDAAEECACEEBAEESCBCSEAD

Этот пример был построен при помощи таблицы случайных чисел ¹⁾.

Б) Пусть при использовании тех же самых пяти букв их вероятности будут 0,4; 0,1; 0,2; 0,2; 0,1 соответственно, причем буквы выбираются независимо. Типичным сообщением такого источника является тогда:

AAACDCBDCEAADADACEDAEADCABEDADDCECAAAAAD

В) Более сложная структура получается, если последовательные символы выбираются не независимо, а так, что их вероятности зависят от предыдущих букв. В простейшем случае такого типа выбор зависит только от предшествующей буквы, но не от ранее стоящих букв. Тогда статистическая структура может быть описана с помощью множества переходных вероятностей $p_i(j)$, т. е. вероятностей того, что за буквой i следует буква j . Индексы i и j пробегают значения, соответствующие всем возможным символам. Другой эквивалентный способ задания этой структуры заключается в том, чтобы задать вероятности «диграмм» (двухбуквенных сочетаний) $p(i,j)$, т. е. относительные частоты диграмм ij . Частоты букв $p(i)$ (вероятности буквы i), переходные вероятности $p_i(j)$ и веро-

¹⁾ Kendall M. and Smith J., Tables of random sampling number, Cambridge, 1939.

$$p(i) = \sum_j p(i, j) = \sum_j p(j, i) = \sum_j p(j) p_j(i),$$

$$\sum_j p_i(j) = \sum_i p(i) = \sum_{i,j} p(i,j) = 1.$$

$p_i(j)$	j A B C	i	$p(i)$	$p(i,j)$	j A B C
A	0 $\frac{4}{5}$ $\frac{1}{5}$	A	$\frac{9}{27}$	A	0 $\frac{4}{15}$ $\frac{1}{15}$
B	$\frac{1}{2}$ $\frac{1}{2}$ 0	B	$\frac{16}{27}$	B	$\frac{8}{27}$ $\frac{8}{27}$ 0
C	$\frac{1}{2}$ $\frac{2}{5}$ $\frac{1}{10}$	C	$\frac{2}{27}$	C	$\frac{1}{27}$ $\frac{4}{135}$ $\frac{0}{135}$

ABBABABABABABBBABBBBBBABABABABBBACACABB
ABBBBABABACBBBABA

Г) Можно также определять вероятностные процессы, которые создают текст, состоящий из последовательности «слов». Предположим, что в языке имеются пять букв А, В, С, D, Е и 16 «слов» с вероятностями появления:

0,10 A	0,16 BEBE	0,11 CABED	0,04 DEB
0,04 ADEB	0,04 BED	0,05 CEED	0,15 DEED
0,05 ADEE	0,02 BEED	0,08 DAB	0,01 EAB
0,01 BADD	0,05 CA	0,04 DAD	0,05 EE

Предположим, что последовательные «слова» выбираются независимо и отделяются друг от друга некоторым промежутком. Типичное сообщение могло бы быть таким:

DAB EE A BEBE DEED DEB ADEE ADEE EE DEB BEBE
BEBE BEBE ADEE BED DEED DEED CEED ADEE A DEED
DEED BEBE CABED BEBE BED DAB DEED ADEB

Если все слова конечной длины, то этот процесс эквивалентен процессу предыдущего типа, но описание может быть проще в терминах структуры слов и их вероятностей. Можно также обобщить и этот случай, вводя переходные вероятности между словами и т. д.

Такие искусственные языки полезны при конструировании простых задач и примеров, имеющих иллюстративные цели. С помощью ряда простых искусственных языков можно также приблизиться к естественному языку. Приближение нулевого порядка получится, если выбирать все буквы с одинаковой вероятностью и независимо. Приближение первого порядка получится, если последовательные буквы выбираются независимо, но каждая буква при этом имеет ту же самую вероятность, что и в естественном языке¹⁾. Поэтому в приближении первого порядка к английскому языку буква E выбирается с вероятностью 0,12 (ее частота в нормативном английском языке) и W с вероятностью 0,02, но нет никакой зависимости между смежными буквами и никакой тенденции образовывать предпочтительные диграммы, такие, как TH, ED и т. д. Для приближения второго порядка вводится структура диграмм. После того как выбрана некоторая буква, следующая буква выбирается в соответствии с частотами, с которыми различные буквы следуют за первой буквой. Для этого требуется таблица частот диграмм $p_i(j)$. Для приближения третьего порядка вводится структура триграмм. Каждая буква выбирается с вероятностями, которые зависят от двух предыдущих букв.

3. Последовательные приближения к английскому языку

Чтобы дать наглядную картину того, как эти последовательные приближения аппроксимируют естественный язык, ниже приводятся типичные последовательности букв для таких приближений к английскому языку. Во всех случаях использовался 27-буквенный «алфавит» (26 букв и пробел между буквами).

¹⁾ Частоты букв, диграмм и триграмм даны в работе: Pratt F., Secret and urgent, Blue Ribbon Book, 1939. Частоты слов табулированы в работе: Dewey G., Relative frequency of english speech sounds, Harvard University Press, 1923.

1. Приближение нулевого порядка (символы независимы и равновероятны):

XFOML RXKHRJFFJUJ ZLPWCFWKCYJ FFJEYVKCQSGHYD
QPAAMKBZAACIBZLHJQD.

2. Приближение первого порядка (символы независимы, но с частотами, свойственными английскому тексту):

OCRO HLI RGWR NMIELWIS EU LL NBNESEBYA TH EEI
ALHENTTPA OOBTTVA NAN BRL.

3. Приближение второго порядка (структура диграмм такая же, как в английском языке):

ON IE ANTSOUTINYS ARE T INCTORE ST BE S DEAMY ACHIN
D ILONASIVE TUCOOME AT TEASONARE FUSO TIZIN ANDY
TOBE SEACE CTISBE.

4. Приближение третьего порядка (структура триграмм такая же, как в английском языке).

IN NO IST LAT WHEY CRATICT FROURE BIRS GROSID
PONDENOME OF DEMONSTURES OF THE REPTAGIN IS
REGOACTIONA OF CRE.

5. Приближение первого порядка на уровне слов. Вместо того чтобы продолжать процесс приближения с помощью структур тетраграмм, . . . , n -грамм, легче и лучше сразу перейти к словарным единицам. Здесь слова выбираются независимо, но с соответствующими им частотами:

REPRESENTING AND SPEEDILY IS AN GOOD APT OR COME
CAN DIFFERENT NATURAL HERE HE THE A IN CAME THE
TO OF TO EXPERT GRAY COME TO FURNISHES THE LINE
MESSAGE HAD BE THESE.

6. Приближение второго порядка на уровне слов. Переходные вероятности от слова к слову являются правильными, но никакая дальнейшая структура не учитывается:

THE HEAD AND IN FRONTAL ATTACK ON AN ENGLISH
WRITER THAT THE CHARACTER OF THIS POINT IS THE
REFORE ANOTHER METHOD FOR THE LETTERS THAT THE
TIME OF WHO EVER TOLD THE PROBLEM FOR AN UNEX-
PECTED¹⁾.

¹⁾ Приведем аналогичным образом построенные примеры для русского языка:

1. ФЮНАЩРЪФЪНЩЦЖЫКАПМЪННЯПЩМНЖЮЧГПМ ЮЮВСТШ-
ЖЕЩЮКЯПЛЧНЩШФОМЕЦЕЭДФБКТТР МЮЕТ

2. ИВЯЫДТАОАДПИ САНЫАЦУЯСДУДЯЪЛЛЯ Л ПРЕБЕ БАЕ-
ОВД ХНЕ АОЛЕТЛС И.

С каждым из шагов, проделанных выше, сходство с обычным английским текстом возрастает довольно заметно. Отметим, что эти примеры имеют достаточно хорошую структуру в пределах расстояний, которые приблизительно в два раза превышают расстояния, учтенные при конструировании. Например, в случае 3 статистический процесс обеспечивает формирование приемлемого текста для двухбуквенных последовательностей, но и четырехбуквенные последовательности из этой выборки обычно могут быть вставлены в осмысленные предложения. В примере 6 последовательности из четырех или более слов могут быть довольно легко вставлены в предложения без необычных или натянутых конструкций. Последовательность из десяти слов «attack on an English writer that the character of this» не является совершенно неприемлемой. Таким образом, оказывается, что достаточно сложный вероятностный процесс дает удовлетворительное представление дискретного источника.

Первые два примера были построены с помощью таблиц случайных чисел, а также (для примера 2) таблицы частот различных букв.

Точно так же можно было бы построить и примеры 3, 4 и 5, так как частоты диграмм, триграмм и отдельных слов известны, но мы использовали более простой эквивалентный метод. Например, чтобы построить пример 3, можно открыть книгу случайным образом и выбрать также случайно букву на странице. Эта буква записывается. Затем книга открывается на другой странице и читается до тех пор, пока не встречается записанная буква. Следующая за ней буква записывается. Затем на другой странице ищется эта последняя буква и записывается следующая за ней и так далее. Аналогичный процесс был использован для составления примеров 4, 5 и 6. Было бы интересно сделать дальнейшие приближения, но на следующей стадии необходимая для этого работа становится огромной.

3. ОТЕ ДОСТОРО ННЕДИЯРИТРКИЯ ПРНОПРОСЕБЫ ЁРЕТ ОС-
КАЛАСИВИ ОМ Р ВШЕРГУ П.

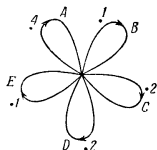
4. С ВОЗДРУНИТЕЛЫБКОТОРОЧЕНЯЛ МЕСЛОСТОЧЕМ МИ ДО.

5. СВОБОДОЙ ДУШЕ ПРОТЯНУЛ КАК ГОВОРИТ ВСПОМНИТЬ
МИЛОСТЬ КОМНАТАМ РАССКАЗА ЖЕНЩИНЫ МНЕ ТУДА ПОНЮ-
ХАВШЕГО КОНЦУ ИСКУСНО КАЖДОМУ РЯСАХ К ДРУГ ПЕРЕРЕЗА-
ЛО ВИДНО ВСЕМ НАЧИНАЕТЕ НАД ДВУХ ЭТО СВЕТА ХОДУНОМ
ЗЕЛЕНАЯ МУХА ЗВУК ОН БЫ ШЕЮ УТЕР БЕЗДАРНЫХ.

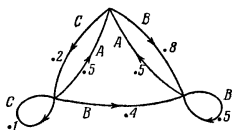
6. ОБЩЕСТВО ИМЕЛО ВЫРАЖЕНИЕ МГНОВЕННОГО ОРУДИЯ К
ДОСТИЖЕНИЮ ДОЛЖНОСТЕЙ ОДИН В РАСЧЕТЫ НА БЕЗНРАВСТ-
ВЕННОСТИ В ПОЭЗИИ РЕЗВИТЬСЯ ВСЕ ГРЫЗЕТ СВОИ БРАЗДЫ
ПРАВЛЕНИЯ НАЧАЛА ЕГО ПОШЛОЙ. — *Прим. ред.*

4. Графическое представление марковского процесса

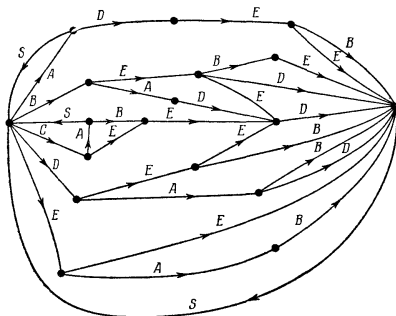
Вероятностные процессы описанного выше типа известны в математике как дискретные марковские процессы, которые подробно



Р и с. 3. Граф, соответствующий источнику в примере Б.



Р и с. 4. Граф, соответствующий источнику в примере В.



Р и с. 5. Граф, соответствующий источнику в примере Г.

изучены и описаны в литературе¹⁾. Общий случай может быть описан следующим образом: существует конечное число возможных «состояний» системы: $S_1, \dots, S_n$. Кроме того, имеется совокупность пере-

¹⁾ Детальное изложение см. Frechet M., *Methods des fonction arbitraires. Théorie ininements en chaine dans le cas d'un nombre fini d'états possibles*, Paris, Cauihier Villars, 1938 [см. также Феллер В., Введение в теорию вероятностей и ее приложения, ИЛ, М., 1952, гл. XV, и Сарымсаков Т. А., Основы теории цепей Маркова, Гостехиздат, М., 1954.— Прим. ред.]

ходных вероятностей $p_i(j)$, т. е. вероятностей того, что система, находящаяся в состоянии S_i , перейдет затем в состояние S_j . Чтобы использовать этот марковский процесс в качестве источника сообщений, нужно только предположить, что при каждом переходе из одного состояния в другое создается одна буква. Состояния будут соответствовать «остатку влияния» предшествовавших букв.

Это положение может быть графически проиллюстрировано, как показано на рис. 3, 4 и 5. «Состояниями» являются узловые точки схемы, а переходные вероятности и создаваемые при этом буквы указаны около соответствующих линий. Рис. 3 соответствует примеру Б в разделе 2, рис. 4 — примеру В. На рис. 3 имеется только одно состояние, так как последовательные буквы независимы. На рис. 4 имеется столько же состояний, сколько букв. Если бы конструировался триграммный пример, то имелось бы самое большее n^2 состояний, соответствующих возможным парам букв, предшествующих выбираемой букве. Рис. 5 представляет схему для случая структуры слов в примере Г. Здесь S соответствует символу «пробел между словами».

5. Эргодические и смешанные источники

Как указывалось выше, для наших целей можно считать, что дискретный источник представляется некоторым марковским процессом. Среди всевозможных дискретных марковских процессов имеется одна группа процессов с особо важными для теории связи свойствами. Этот специальный класс состоит из «эргодических» процессов, и мы будем называть соответствующие источники эргодическими источниками. Хотя строгое определение эргодического процесса несколько сложно, общая идея проста. В случае эргодического процесса каждая создаваемая процессом последовательность имеет одни и те же статистические свойства. Поэтому частоты букв, диграмм и т. п., полученные из частных последовательностей, будут стремиться с увеличением длин последовательностей к определенным пределам, не зависящим от этих частных последовательностей. В действительности это верно не для каждой последовательности, но множество последовательностей, для которых это не верно, имеет вероятность, равную нулю. Грубо говоря, свойство эргодичности означает статистическую однородность.

Все примеры искусственных языков, данные выше, являются эргодическими. Это свойство связано со структурой соответствующей схемы. Процесс будет эргодическим, если соответствующий граф обладает следующими двумя свойствами¹⁾.

¹⁾ Эти условия являются переформулированными в терминах графов условиями, данными в книгах, указанных в примечании на стр. 256; также даны в работах Феллера и Сарымсакова, указанных в приложенной к сборнику библиографии. — *Прим. ред.*

1. Соответствующая схема не распадается на две изолированные части A и B , такие, что от одной узловой точки в части A нельзя было бы перейти в направлении стрелок в точки части B , и наоборот.

2. Каждая замкнутая последовательность линий, стрелки которых ориентированы в одном направлении, называется «циклом». Под «длиной цикла» понимается число линий, из которых он состоит. Так, например, на рис. 5 последовательность BEBES есть цикл длины 5. Второе свойство состоит в том, чтобы наибольший общий делитель длин всех циклов равнялся единице¹⁾.

Если первое условие удовлетворено, а второе нарушено тем, что общий делитель $d > 1$, то последовательности имеют некоторого рода периодическую структуру. Различные последовательности распадаются на d различных классов, которые в статистическом отношении одинаковы, за исключением сдвига начала (т. е. выбора того, какую букву последовательности назвать первой). С помощью смещения на величину от 0 до $d-1$ каждая последовательность может быть сделана статистически эквивалентной любой другой. При $d = 2$ простым примером является следующий: имеются три возможные буквы a, b, c . За буквой a следует либо b , либо c с вероятностями $1/3$ и $2/3$ соответственно. За b и c всегда следует буква a . Тогда типичная последовательность имеет вид:

abacacacabacababacac

Процессы такого типа не будут иметь большого значения для нашей работы.

Если нарушено первое условие, то граф может быть разделен на некоторое число подграфов, каждый из которых удовлетворяет первому условию. Будем предполагать, что второе условие также выполняется для каждого подграфа. В этом случае имеет место то, что может быть названо «смешанным» источником, составленным из некоторого числа чистых компонент²⁾. Эти компоненты соответствуют различным подграфам. Если $L_1, L_2, L_3, \dots$ — источники, соответствующие этим компонентам, то можно записать

$$L = p_1 L_1 + p_2 L_2 + p_3 L_3 + \dots,$$

где p_i — вероятность компоненты L_i ³⁾.

Физический смысл описанного состоит в следующем: имеется несколько различных источников $L_1, L_2, L_3, \dots$, каждый из которых имеет однородную статистическую структуру (т. е. является эрго-

¹⁾ В математической литературе принято называть эргодическими процессы, обладающие свойством 1) и, быть может, не обладающие свойством 2). Как отмечает далее автор, различие между этими двумя классами процессов малосущественно для дальнейшего изложения. — *Прим. ред.*

²⁾ В математической литературе эти компоненты называют классами состояний. — *Прим. ред.*

³⁾ Это равенство имеет, конечно, лишь условный смысл. — *Прим. ред.*

дическим). Априори неизвестно, какой источник будет использован, но если последовательность начинается с состояния, принадлежащего данной чистой компоненте L_i , то она продолжается бесконечно в соответствии со статистической структурой этой компоненты.

В качестве примера можно взять два процесса из определенных выше и предположить, что $p_1 = 0,2$ и $p_2 = 0,8$. Последовательность из смешанного источника

$$L = 0,2L_1 + 0,8L_2$$

могла бы быть получена следующим образом: сначала выбирается L_1 или L_2 с вероятностями 0,2 и 0,8, а затем выбранный источник создает последовательность.

Если не оговорено противное, будем предполагать, что источник является эргодическим. Такое предположение позволяет отождествлять средние значения вдоль некоторой последовательности со средними значениями по ансамблю возможных последовательностей (причем вероятность расхождения равна нулю). Например, относительная частота буквы А в частной бесконечной последовательности будет с вероятностью единица равняться ее относительной частоте по ансамблю последовательностей.

Если p_i — вероятность состояния i , а $p_i(j)$ — вероятность перехода в состояние j , то для того, чтобы процесс был стационарным, p_i должны, очевидно, удовлетворять условиям равновесия:

$$p_j = \sum_i p_i p_i(j).$$

Можно показать, что в эргодического случае при любых начальных условиях вероятности пребывания в состоянии j после N шагов $p_j(N)$ при $N \rightarrow \infty$ стремятся к величинам, удовлетворяющим условиям равновесия.

6. Выбор, неопределенность и энтропия

Дискретный источник информации был представлен как марковский процесс. Можно ли определить величину, которая будет измерять в некотором смысле, как много информации создается таким процессом, или, лучше, с какой скоростью она создается?

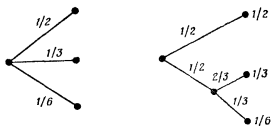
Предположим, что имеется некоторое множество возможных событий, вероятности осуществления которых суть $p_1, p_2, \dots, p_n$. Эти вероятности известны, но это — все, что нам известно относительно того, какое событие произойдет. Можно ли найти меру того, насколько велик «выбор» из такого набора событий или сколь неопределенен для нас его исход?

Если имеется такая мера, скажем $H(p_1, p_2, \dots, p_n)$, то разумно потребовать, чтобы она обладала следующими свойствами:

1. H должна быть непрерывной относительно p_i .

2. Если все p_i равны, $p_i = \frac{1}{n}$, то H должна быть монотонно возрастающей функцией от n . В случае равновероятных событий имеется больше возможностей выбора или неопределенности, чем в случае, когда имеются разноравные события.

3. Если бы выбор распадался на два последовательных выбора, то первоначальная H должна была бы быть взвешенной суммой индивидуальных значений H . Смысл этого иллюстрируется



Р и с. 6. Выбор из трех возможностей.

рис. 6. Слева имеются три возможности $p_1 = 1/2$; $p_2 = 1/3$; $p_3 = 1/6$. Справа производится выбор между двумя возможностями, причем каждая имеет вероятность $1/2$, и в случае осуществления второй возможности производится еще один выбор между двумя возможностями с вероятностями $2/3$; $1/3$. Окончательные результаты имеют те же самые вероятности, как и прежде. Потребуем в этом конкретном случае, чтобы

$$H\left(\frac{1}{2}, \frac{1}{3}, \frac{1}{6}\right) = H\left(\frac{1}{2}, \frac{1}{2}\right) + \frac{1}{2} H\left(\frac{2}{3}, \frac{1}{3}\right).$$

Коэффициент $1/2$ является весовым множителем, введенным из-за того, что второй выбор осуществляется только в половине всех случаев.

В приложении 2 устанавливается следующее.

Теорема 2. Существует единственная функция H , удовлетворяющая трем перечисленным выше свойствам. При этом H имеет вид

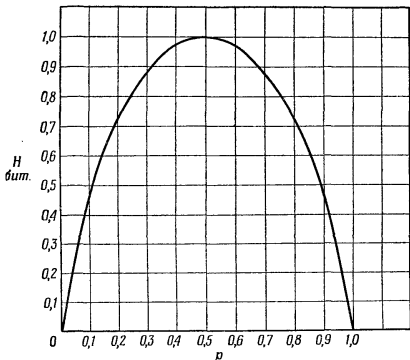
$$H = -K \sum_{i=1}^n p_i \log p_i,$$

где K — некоторая положительная константа.

Эта теорема и допущения, требуемые для ее доказательства, не являются необходимыми для настоящей теории. Они приводятся главным образом с тем, чтобы обосновать целесообразность не-

которых из дальнейших определений. Действительное же оправдание этих определений заключается в том, что из них проистекает¹⁾.

Величины вида $H = - \sum p_i \log p_i$ (постоянная K определяет просто выбор единицы измерения) играют центральную роль в теории информации в качестве меры количества информации,



Р и с. 7. Энтропия в случае двух возможностей с вероятностями p и $(1-p)$.

возможности выбора и неопределенности. Форма величины H оказывается такой же, как и форма энтропии, определяемой в статистической механике²⁾, где p_i — вероятность того, что система находится в ячейке i фазового пространства. Величина H в таком виде встречается, например, в знаменитой теореме Больцмана. Назовем величину $H = - \sum p_i \log p_i$ энтропией множества вероятностей $p_1 \dots p_n$. Если x — случайная величина, то мы обозначим ее энтропию через $H(x)$; таким образом, x — не аргумент

¹⁾ Более подробно доказательство теоремы 2 проводится в работе: Х и н ч и н А. Я., Понятие энтропии в теории вероятностей, *Успехи матем. наук*, VIII, № 3, 3—20. Возможность ослабить систему аксиом 1), 2), 3) отмечена в работах Фаддеева В. К. (Фаддеев В. К., К понятию энтропии конечной вероятностной схемы, *Успехи матем. наук*, 11, (1956), 1.—Прим. ред.,

²⁾ См., например, Т о л м а н R. C., *Principles of statistical mechanics*, Oxford, Clarendon, 1938.

функции, а лишь знак, отличающий ее, скажем, от $H(y)$ — энтропии случайной величины y .

На рис. 7 представлена энтропия для случая двух исходов с вероятностями p и $q = 1 - p$ в виде функции от p , а именно:

$$H = -(p \log p + q \log q).$$

Величина H обладает рядом интересных свойств, которые также подтверждают, что она является разумной количественной мерой возможности выбора или мерой количества информации.

1. $H=0$ тогда и только тогда, когда все вероятности p_i , кроме одной, равны нулю, а эта единственная вероятность равна единице. Таким образом, H равна нулю только в случае полной определенности исхода опыта. В противном случае H положительна.

2. При заданном n величина H максимальна и равна $\log n$, когда все p_i равны (следовательно, $p_i = 1/n$). То, что в этом случае неопределенность будет наибольшей, чувствуется также и интуитивно.

3. Пусть имеются два события x и y с m исходами для первого и n исходами для второго. Пусть $p(i, j)$ означает вероятность совместного осуществления исхода i для x и j для y . Энтропия совместного события равна

$$H(x, y) = - \sum_{i,j} p(i, j) \log p(i, j),$$

в то время как

$$H(x) = - \sum_{i,j} p(i, j) \log \sum_j p(i, j),$$

$$H(y) = - \sum_{i,j} p(i, j) \log \sum_i p(i, j).$$

Легко показать, что

$$H(x, y) \leq H(x) + H(y),$$

причем равенство имеет место только в том случае, когда события независимы [т. е. $p(i, j) = p(i)p(j)$]. Неопределенность совместного события меньше или равна сумме неопределенностей отдельных событий.

4. Всякое изменение вероятностей $p_1, p_2, \dots, p_n$ в сторону их выравнивания увеличивает H . Так, если $p_1 < p_2$ и увеличивать p_1 , уменьшая одновременно p_2 на такую же величину, так что p_1 и p_2 приближаются друг к другу, то H увеличивается. В более общем виде, если над вероятностями p_i произвести операцию «осреднения» вида

$$p'_i = \sum_j a_{ij} p_j,$$

где $\sum_i a_{ij} = \sum_j a_{ij} = 1$ и все $a_{ij} \geq 0$, то H увеличивается (за исключением того частного случая, в котором такое преобразование сводится к одной только перестановке p_j , что, конечно, не изменяет значения H).

5. Пусть имеются два случайных события x и y , как и в п. 3, не обязательно независимые. Для каждого частного значения i , которое может принять x , имеется условная вероятность $p_i(j)$ того, что y при этом примет значение j . Она задается выражением

$$p_i(j) = \frac{p(i, j)}{\sum_j p(i, j)}.$$

Определим *условную энтропию*¹⁾ $H_x(y)$ величины y как величину, получаемую в результате осреднения энтропии y , вычисленной по всем значениям x , с весами, соответственно равными вероятностям этих значений x . Таким образом,

$$H_x(y) = - \sum_{i,j} p(i, j) \log p_i(j).$$

Эта величина показывает, какова в среднем неопределенность значения y , когда известно значение x . Подставляя значение $p_i(j)$, получим

$$\begin{aligned} H_x(y) &= - \sum_{i,j} p(i, j) \log p(i, j) + \sum_{i,j} p(i, j) \log \sum_j p(i, j) = \\ &= H(x, y) - H(x) \end{aligned}$$

или

$$H(x, y) = H(x) + H_x(y).$$

Неопределенность (или энтропия) совместного события (x, y) равна неопределенности события x плюс неопределенность события y , когда x известно.

6. Из п. 3 и 5 имеем

$$H(x) + H(y) \geq H(x, y) = H(x) + H_x(y).$$

Отсюда

$$H(y) \geq H_x(y).$$

Неопределенность события y не возрастает от того, что событие x становится известным. Она уменьшается, если только события x и y не являются независимыми. В противном случае она не изменяется.

¹⁾ В современной математической литературе предпочитают называть условной энтропией сумму $\sum_i p_i(j) \log p_i(j)$, являющуюся функцией от i , а величину $H_x(y)$ называть средней условной энтропией. — Прим. ред.

7. Энтропия источника сообщений

Рассмотрим дискретный источник с конечным числом состояний, наподобие источников, рассмотренных выше. Для каждого возможного состояния i имеется некоторое множество вероятностей $p_i(j)$ создания различных возможных символов j . Следовательно, для каждого состояния i существует энтропия H_i . Энтропия источника определяется как среднее значение величин H_i , которые осреднены в соответствии с вероятностями осуществления соответствующих событий

$$H = \sum_i P_i H_i = - \sum_{i,j} P_i p_i(j) \log p_i(j).$$

Это энтропия источника на символ текста. Если наш марковский процесс развивается с определенной скоростью, то можно говорить также об энтропии в секунду

$$H' = \sum_i f_i H_i,$$

где f_i — средняя частота (появлений в секунду) состояния i . Очевидно,

$$H' = mH,$$

где m — среднее число символов, создаваемых за одну секунду. Величины H или H' измеряют количество информации, создаваемое источником, на символ или в секунду. Если в качестве основания логарифмов выбрано 2, то они представляют из себя количество битов на символ или в секунду.

Если символы в последовательности независимы, то H просто равняется — $\sum p_i \log p_i$, где p_i — вероятность символа i . Предположим, что в этом случае рассматривается длинное сообщение, состоящее из N символов. Оно будет с большой вероятностью содержать $p_1 N$ раз первый символ, $p_2 N$ раз второй символ и т. д. Отсюда вероятность этого конкретного сообщения будет приближенно равна

$$p = p_1^{p_1 N} p_2^{p_2 N} \dots p_n^{p_n N},$$

или

$$\log p = N \sum_i p_i \log p_i,$$

$$\log p = -NH,$$

$$H = -\frac{\log \frac{1}{p}}{N}.$$

Таким образом, H приближенно равна логарифму обратной величины вероятности типичной длинной последовательности, поде-

ленному на число символов в последовательности. Этот вывод верен и для любого источника. В более точной формулировке мы имеем (см. приложение 3) следующую теорему.

Теорема 3. Для любых заданных $\varepsilon > 0$ и $\delta > 0$ можно найти такое N_0 , что последовательности любой длины $N \geq N_0$ распадутся на два класса:

1) множество последовательностей, суммарная вероятность которых меньше чем ε ;

2) статок, все члены которого обладают вероятностями, удовлетворяющими неравенству

$$\left| \frac{\log \frac{1}{p}}{N} - H \right| < \delta.$$

Другими словами, почти достоверно, что $\frac{\log \frac{1}{p}}{N}$ весьма близко к H , когда N велико.

Аналогичный результат получается и для последовательностей, имеющих определенные суммарные вероятности. Рассмотрим снова последовательности длины N . Расположим их в ряд в порядке уменьшения вероятностей. Введем величину $n(q)$ — число последовательностей, которые необходимо взять из нашего ряда, начиная с наиболее вероятной последовательности, для того чтобы для взятых последовательностей накопилась суммарная вероятность q .

Теорема 4. Когда q не равно нулю или единице,

$$\lim_{N \rightarrow \infty} \frac{\log n(q)}{N} = H.$$

Можно интерпретировать величину $\log n(q)$ как число бит, требуемых для задания последовательности, когда рассматриваются только наиболее вероятные последовательности с суммарной вероятностью q . Тогда $\frac{\log n(q)}{N}$ есть число бит на символ, необходимых для задания последовательностей. Теорема гласит, что для больших N это число не зависит от q и равно H . Быстрота возрастания логарифма числа сравнительно вероятных последовательностей определяется величиной H независимо от истолкования смысла слов «сравнительно вероятный». Эти выводы, доказываемые в приложении 3, показывают, что в большинстве случаев длинные последовательности можно рассматривать так, как если бы их было ровно 2^{HN} , и каждая имела вероятность 2^{-HN} .

Следующие две теоремы показывают, что H и H' могут быть найдены с помощью предельных операций непосредственно

из статистической структуры последовательности сообщений без привлечения вероятностей состояний и переходных вероятностей.

Теорема 5. Пусть $p(B_i)$ — вероятность того, что источник создает последовательность символов B_i . Пусть

$$G_N = -\frac{1}{N} \sum_i p(B_i) \log p(B_i),$$

где суммирование распространяется на все последовательности B_i , содержащие N символов. Тогда G_N — монотонно убывающая функция от N и

$$\lim_{N \rightarrow \infty} G_N = H.$$

Теорема 6. Пусть $p(B_i, S_j)$ — вероятность того, что появится последовательность B_i и после нее появится символ S_j , а $p_{B_i}(S_j) = p(B_i, S_j)/p(B_i)$ — условная вероятность того, что после B_i появится S_j . Пусть

$$F_N = -\sum_{i,j} p(B_i, S_j) \log p_{B_i}(S_j),$$

где суммирование проводится по всем последовательностям B_i из $(N-1)$ символа и по всем символам S_j . Тогда F_N — монотонно убывающая функция от N ,

$$F_N = NG_N - (N-1)G_{N-1},$$

$$G_N = \frac{1}{N} \sum_{i=1}^N F_i,$$

$$F_N \leq G_N$$

и

$$\lim_{N \rightarrow \infty} F_N = H.$$

Доказательства этих теорем приводятся в приложении 3. Они показывают, что ряд приближений к H может быть получен с помощью рассмотрения только статистической структуры последовательностей, охватывающей 1, 2, ..., N символов. F_N является лучшим приближением. На самом деле F_N является энтропией N -го приближения к источнику рассмотренного выше типа. Если нет никаких статистических связей, распространяющихся более чем на N символов, т. е. если условная вероятность появления

следующего символа при условии, что значения $(N-1)$ (предшествовавших символов) известно и не изменяется от добавления сведений о любых стоящих ранее символах, то тогда $F_N = H$. F_N является, конечно, условной энтропией следующего символа, когда известны $(N-1)$ (предыдущих), тогда как G_N — энтропия на символ последовательности из N символов.

Отношение энтропии источника к максимальному значению, которого могла бы достичь энтропия при тех же символах, назовем относительной энтропией источника. Как будет показано ниже, она является максимальным сжатием, которое можно осуществить кодированием при помощи того же самого алфавита. Единица минус относительная энтропия есть *избыточность*. Избыточность обычного английского текста, если не рассматривать статистическую структуру, относящуюся более чем к 8 буквам, составляет примерно 50%. Это значит, что когда мы пишем по-английски, то половина знаков текста определяется структурой языка, и лишь половина выбирается свободно. Значение 50% было найдено с помощью нескольких независимых методов, которые все дают приблизительно тот же самый результат. Одним из таких методов было вычисление энтропии приближений к английскому тексту. Другой метод состоял в исключении из образцов английского текста некоторой части букв, после чего делалась попытка их восстановить. Если их удастся восстановить, когда 50% текста исключено, избыточность должна быть больше чем 50%. Третий метод связан с некоторыми известными данными работ по криптографии.

Два крайних примера избыточности английской прозы представляют из себя бейсик-инглиш¹⁾ и книга Джемса Джойса²⁾. Словарь языка бейсик-инглиш ограничен 850 словами и избыточность его очень велика. Это выражается в увеличении объема текста, которое происходит при переводе на бейсик-инглиш. С другой стороны, Джойс расширяет словарь и это позволяет ему достигнуть сжатия смыслового содержания.

Избыточность языка связана с возможностью построения кроссвордов. Если избыточность равна нулю, то любая последовательность букв представляет собой осмысленный текст, и любое двумерное построение букв представляет из себя кроссворд. Если избыточность слишком велика, то язык налагает столь много ограничений, что уже не могут существовать большие кроссворды.

Более точный анализ показывает, что если предположить случайный характер ограничений, налагаемых языком, то большие кросс-

¹⁾ Бейсик-инглиш — это весьма упрощенный вариант английского языка. — *Прим. ред.*

²⁾ J o u s e J a m e s, *Finegans Wake*.

ворды становятся возможными, лишь когда избыточность достигает 50%. Если избыточность достигает 33%, то возможны трехмерные кроссворды и т. д.

8. Представление операций кодирования и декодирования

Нам нужно теперь представить математически операции, выполняемые передатчиком и приемником при кодировании и декодировании информации. И передатчик и приемник будем называть дискретными преобразователями. На вход преобразователя поступает последовательность входных символов, а на выходе получается последовательность выходных символов. Преобразователь может обладать внутренней памятью, так что выходной символ зависит не только от данного входного символа, но и от всех предыдущих. Предположим, что внутренняя память конечна, т. е. существует конечное число m возможных состояний преобразователя и при этом очередной выходной символ является функцией от соответствующего входного символа и от состояния, в котором находится преобразователь. Следующее состояние преобразователя будет другой функцией от этих же аргументов. Таким образом, преобразователь может быть описан двумя функциями:

$$y_n = f(x_n, \alpha_n),$$

$$\alpha_{n+1} = g(x_n, \alpha_n),$$

где x_n — n -й входной символ,

α_n — состояние преобразователя к моменту ввода n -го входного символа.

y_n — выходной символ (или группа выходных символов), создаваемый, когда на вход поступает x_n и преобразователь находится в состоянии α_n .

Если выходные символы одного преобразователя могут служить входными символами для другого, то преобразователи могут быть соединены последовательно, в результате чего также получится некоторый преобразователь. Если существует второй преобразователь, который при соединении его входа с выходом первого восстанавливает исходный входной сигнал, то первый преобразователь называется невырожденным, а второй — обратным к первому.

Теорема 7. Выход преобразователя с конечным числом состояний, подключенного к статистическому источнику с конечным числом состояний, представляет из себя статистический источник с конечным числом состояний, причем энтропия этого источника (в единицу времени) меньше или равна энтропии источника на входе. Если преобразователь невырожденный, то энтропии равны.

Пусть α представляет собой состояние источника, который создает последовательность символов x_i , и пусть β — состояние преобразователя, который создает на выходе блоки символов y_i . Система, полученная в результате соединения источника с преобразователем, может быть представлена с помощью произведения «пространств состояний», состоящего из пар (α, β) . Две точки в этом пространстве (α_1, β_1) и (α_2, β_2) соединяются линией, если источник может изменить состояние α_1 на состояние α_2 , создав при этом такой символ x , который изменит состояние преобразователя β_1 на β_2 . Этой линии приписывается вероятность символа x , и она помечается блоком символов — y_1 , который был создан преобразователем при переходе из состояния β_1 в состояние β_2 . Энтропия источника на выходе может быть вычислена как взвешенная сумма по всем состояниям. Если мы просуммируем сначала по β , то каждый получающийся в результате член будет не больше соответствующего члена для α , и, следовательно, энтропия не увеличивается. Если преобразователь невырожденный, то к его выходу можно присоединить второй преобразователь, обратный первому. Если при этом H'_1 , H'_2 и H'_3 представляют соответственно энтропии источника и выходов первого и второго преобразователей, то $H'_1 \geq H'_2 \geq H'_3 = H'_1$ и, следовательно, $H'_1 = H'_2$.

Пусть на возможные последовательности наложена система ограничений того типа, который может быть представлен графически, как на рис. 2. Если бы различным линиям, соединяющим состояние i с состоянием j , были приписаны вероятности $p_{ij}^{(s)}$, то эта система стала бы источником. Имеется один частный способ приписывания вероятностей, при котором энтропия становится максимальной (см. приложение 4).

Теорема 8. Пусть система ограничений, рассматриваемая как канал, имеет пропускную способность $C = \log W$. Если положить

$$p_{ij}^{(s)} = \frac{B_j}{B_i} \cdot W^{-l_{ij}^{(s)}},$$

где $l_{ij}^{(s)}$ — длительность s -го символа, ведущего от состояния i к состоянию j , а B_i удовлетворяет условию

$$B_i = \sum_{s,j} B_j W^{-l_{ij}^{(s)}},$$

то энтропия H максимальна и равна C .

С помощью надлежащего выбора величин переходных вероятностей энтропия символов на выходе канала может быть доведена до максимума, являющегося пропускной способностью канала.

9. Основная теорема для канала без шума¹⁾

Подтвердим теперь нашу интерпретацию величин H как скорости создания информации доказательством того факта, что H определяет пропускную способность канала, необходимую при наиболее эффективном кодировании.

Теорема 9. Пусть источник имеет энтропию H (бит на символ), а канал имеет пропускную способность C (бит в секунду). Тогда можно закодировать сообщения на выходе источника таким образом, чтобы передавать символы по каналу со средней скоростью $C/H - \epsilon$ символов в одну секунду, где ϵ — сколь угодно мало. Передавать со средней скоростью, большей чем C/H , невозможно.

Обратная часть теоремы, утверждающая, что нельзя превзойти скорость C/H , может быть доказана, если заметить, что энтропия в секунду на входе канала равна энтропии источника, так как передатчик должен быть невырожденным преобразователем, и что эта энтропия не может превосходить пропускной способности канала. Отсюда $H' \leq C$ и число символов в секунду равно $H'/H \leq C/H$.

Докажем первую часть теоремы двумя различными способами. Первый способ состоит в рассмотрении множества всех последовательностей из N символов, создаваемых источником. При большом N все эти последовательности можно разделить на две группы, одна из которых содержит меньше чем $2^{(H+\eta)N}$ членов, а вторая содержит меньше чем 2^{RN} членов (где R — логарифм числа различных символов) и имеет суммарную вероятность, меньшую чем μ . С ростом N величины η и μ стремятся к нулю. Число сигналов в канале, имеющих длительность T , больше чем $2^{(C-\theta)T}$, где θ мало, когда T велико. Если выбрать

$$T = \left(\frac{H}{C} + \lambda \right) N,$$

то найдется достаточное число последовательностей символов в канале для высоковероятностной группы, когда N и T достаточно велики (сколь бы малым ни было выбрано λ), и несколько добавочных последовательностей. Высоковероятностная группа произвольным, взаимнооднозначным образом кодируется в это множество последовательностей символов канала. Остающиеся последовательности источника кодируются в более длинные последовательности канала, начинающиеся и заканчивающиеся одной из добавочных последовательностей, не использованных для высоковероятностной группы, причем эта последовательность действует как начальный и конечный сигналы, указывающие на использование в промежутке иного кода. Между эти-

¹⁾ Подробное изложение части результатов этого раздела можно найти в работе Хинчина; см. сноску 1, стр. 261. — Прим. ред.

ми сигналами оставляют временной интервал, достаточный для того, чтобы для этого временного интервала в канале существовало достаточно различных последовательностей для всех маловероятных сообщений. Для этого потребуется, чтобы такой временной интервал равнялся

$$T_1 = \left(\frac{R}{C} + \varphi \right) N,$$

где φ мало. Средняя скорость передачи символов сообщения в 1 секунду будет тогда больше чем

$$\left[(1-\delta) \frac{T}{N} + \delta \frac{T_1}{N} \right]^{-1} = \left[(1-\delta) \left(\frac{H}{C} + \lambda \right) + \delta \left(\frac{R}{C} + \varphi \right) \right]^{-1}.$$

При увеличении N величины δ , λ и φ стремятся к нулю и скорость стремится к C/H .

Другой способ выполнения такого кодирования и, следовательно, другой способ доказательства теоремы можно описать следующим образом: расположим сообщения длины N в порядке убывания их вероятностей; пусть эти вероятности будут

$p_1 \geq p_2 \geq \dots \geq p_n$. Пусть $P_s = \sum_{i=1}^{s-1} p_i$, т. е. P_s — накопленная

вероятность до p_{s-1} включительно. Закодируем сначала все сообщения в двоичную систему. Двоичный код для сообщения S получается путем разложения P_s как двоичного числа. Разложение проводится до m_s -й позиции, где m_s — целое число, удовлетворяющее соотношению

$$\log_2 \frac{1}{p_s} \leq m_s < 1 + \log_2 \frac{1}{p_s}.$$

Таким образом, высоковероятные сообщения представляются короткими кодами, а маловероятные — длинными. Из этих неравенств вытекает следующее неравенство:

$$\frac{1}{2^{m_s}} \leq p_s < \frac{1}{2^{m_s-1}}.$$

Код для P_s будет отличаться от всех последующих кодов одной или более из своих m_s позиций, так как все остающиеся P_i по крайней мере на 2^{-m_s} больше и поэтому их двоичные разложения отличаются от кода для P_s на первых m_s позициях. Следовательно, все эти коды различны, и можно восстановить сообщение по его коду. Если последовательности символов канала не являются уже двоичными последовательностями, им можно приписать двоичные числа произвольным образом и преобразовать таким образом двоичный код в сигналы, используемые для канала.

Среднее число H_1 бит, употребляемых на символ первоначального сообщения, легко оценить. Имеем

$$H_1 = \frac{1}{N} \sum m_s p_s.$$

Но

$$\frac{1}{N} \sum \left(\log_2 \frac{1}{p_s} \right) p_s \leq \frac{1}{N} \sum m_s p_s < \frac{1}{N} \sum \left(1 + \log_2 \frac{1}{p_s} \right) p_s$$

и поэтому

$$G_N \leq H_1 < G_N + \frac{1}{N}.$$

С ростом N величина G_N стремится к H — энтропии источника, и H_1 также стремится к H .

Отсюда видно, что неэффективность кодирования в случае, когда используется конечное запаздывание на N символов, не должна быть больше, чем $1/N$ плюс разность между истинной энтропией H и энтропией G_N , сосчитанной для последовательностей длины N . Поэтому необходимое увеличение времени по сравнению с идеальным в процентах не превысит

$$\frac{G_N}{H} + \frac{1}{HN} - 1.$$

Этот метод кодирования в сущности совпадает с методом, независимо найденным Р. М. Фано¹⁾. Его метод состоит в расположении сообщений длины N в порядке убывающих вероятностей. Этот ряд делится на две группы, по возможности с равными вероятностями. Если сообщение относится к первой группе, его первая двоичная цифра будет 0, в противном случае — единица. Эти группы аналогичным образом делятся на подгруппы примерно равной вероятности, и частная подгруппа определяет второй двоичный знак. Этот процесс продолжается до тех пор, пока не получатся подгруппы, содержащие только по одному сообщению. Легко видеть, что, за исключением незначительных отличий (в общем случае в последней цифре), это приводит к тем же результатам, что и при описанном выше арифметическом методе.

10. Рассмотрение полученных результатов и примеры

Для передачи максимальной мощности от генератора в нагрузку в общем случае применяется трансформатор, делающий сопротивление генератора, наблюдаемое со стороны нагрузки, равным сопротивлению этой нагрузки. Это аналогично рассмотренной выше ситуации: преобразователь, осуществляющий кодирование, согласовывает источник с каналом в статистическом смысле. Источник, рассматриваемый через преобразователь со стороны канала, должен иметь такую же статистическую структуру, какую имеет источник,

¹⁾ Fano R. M., Technical № 65, The Research Laboratory of Electronics, MIT, March 17, 1949.

который максимизирует энтропию в канале. Содержание теоремы 9 сводится к тому, что, хотя точное согласование в общем случае невозможно, к нему можно приблизиться сколь угодно близко. Отношение действительной скорости передачи к пропускной способности C можно назвать эффективностью системы кодирования. Оно равно отношению действительной энтропии символов канала к максимально возможной энтропии.

Вообще говоря, идеальное или приблизительно идеальное кодирование требует длительных задержек в передатчике и приемнике. В случае отсутствия шума в канале, который мы сейчас рассмотрели, главное назначение этой задержки состоит в том, что она позволяет достаточно хорошо согласовать вероятности с соответствующими длительностями последовательностей. Для хорошего кода логарифм обратной величины вероятности длинного сообщения должен быть пропорционален длительности соответствующего сигнала; в действительности величина

$$\left| \frac{\log \frac{1}{p}}{T} - C \right|$$

должна быть мала для всех длинных сообщений, за исключением небольшой их доли.

Если источник может создавать только одно определенное сообщение, то его энтропия равна нулю и никакой канал не нужен. Например, счетная машина, спроектированная для вычисления десятичного разложения числа π , производит определенную последовательность без всяких элементов случайности. Для «передачи» этой последовательности в другой пункт не требуется никакого канала. В этом пункте можно построить другую машину, вычисляющую ту же самую последовательность. Однако это может быть непрактично. В таком случае можно пренебречь некоторыми или всеми статистическими характеристиками, которые известны об источнике. Можно принять, что цифры числа π образуют случайную последовательность, и сконструировать систему, способную передавать любую последовательность цифр. Точно так же при конструировании кода можно использовать лишь некоторые из известных статистических характеристик английского языка, а не все.

В этом случае рассматривается источник с максимальной энтропией, подчиненный статистическим условиям, которые мы пожелаем сохранить. Энтропия этого источника определяет необходимую и достаточную пропускную способность канала. В примере с числом π оставлены только те сведения, что все цифры выбираются из множества $0, 1, \dots, 9$. В примере с английским языком можно было бы пожелать достичь экономии за счет статистики, используя лишь знания частот букв. Источник с мак-

симальной энтропией будет тогда первым приближением к английскому языку и его энтропия определит необходимую пропускную способность канала.

11. Примеры

В качестве простого примера использования некоторых из полученных результатов рассмотрим источник, создающий последовательность букв, выбранных из набора букв A, B, C, D с вероятностями $1/2, 1/4, 1/8, 1/8$, причем очередная буква выбирается независимо. Имеем

$$H = - \left(\frac{1}{2} \log \frac{1}{2} + \frac{1}{4} \log \frac{1}{4} + \frac{2}{8} \log \frac{1}{8} \right) = \frac{7}{4} \text{ бит/символ.}$$

Таким образом, для кодирования сообщений этого источника двоичными знаками в пределе достаточно в среднем $7/4$ битов на символ. В этом случае можно действительно достигнуть предельного значения, применяя следующий код (полученный по методу второго доказательства теоремы 9):

A	0
B	10
C	110
D	111

Среднее число битов, потребных для кодирования последовательности из N символов, будет равно

$$N \cdot \left(\frac{1}{2} \times 1 + \frac{1}{4} \times 2 + \frac{2}{8} \times 3 \right) = \frac{7}{4} N.$$

Легко видеть, что двоичные знаки 0 и 1 имеют вероятности $1/2, 1/2$, так что H для кодированных последовательностей равна одному биту на символ. Так как в среднем мы имеем $7/4$ двоичных знаков на первоначальную букву, то энтропия на единицу времени будет той же самой. Максимальная возможная энтропия для первоначального набора букв равна $\log_2 4 = 2$ и достигается, когда A, B, C, D обладают вероятностями $1/4, 1/4; 1/4; 1/4$. Отсюда относительная энтропия равна $7/8$. Можно перевести двоичные последовательности в первоначальные символы в соотношении 2 к 1, основываясь на следующей таблице:

00	A'
01	B'
10	C'
11	D'

Тогда этот двойной процесс закодирует первоначальное сообщение в те же самые символы, но со средним коэффициентом сжатия $7/8$.

В качестве второго примера рассмотрим источник, который создает последовательность из букв A и B , выбираемых с вероятностями p для A и q для B . Если $p \ll q$, то имеем

$$H = -\log_2 p^p \cdot (1-p)^{1-p} = -p \log_2 p (1-p)^{\frac{1-p}{p}} = p \log_2 \frac{e}{p}.$$

В этом случае можно создать достаточно хорошую систему кодирования сообщений для двоичного канала, посылая специальный символ, скажем 0000, в случае, когда нужно передать редкий символ A , а затем, посылая последовательность, указывающую число букв B , следующих за ним. Это число может быть указано путем представления в двоичной системе, причем все числа, содержащие специальную последовательность, исключаются. Все числа до 16 изображаются, как обычно; 16 представляется следующим после 16 двоичным числом, которое не содержит четырех нулей, а именно $17 = 10001$ и т. д.

Можно показать, что при $p \rightarrow 0$ это кодирование приближается к идеальному, если только длина специальной последовательности выбрана правильно.

II. ДИСКРЕТНЫЙ КАНАЛ С ШУМОМ

11. Представление дискретного канала с шумом

Рассмотрим теперь случай, когда в процессе передачи сигнал искажается шумом. Это означает, что принятый сигнал не обязательно совпадает с сигналом, посланным передатчиком. Можно различать два случая. Если из определенного переданного сигнала всегда получается один и тот же принятый сигнал, т. е. принятый сигнал является определенной функцией от переданного сигнала, то такое явление может быть названо искажением. Если эта функция имеет обратную, т. е. никакие два переданных сигнала не создают один и тот же принятый сигнал, то искажение может быть скорректировано, по крайней мере в принципе, просто путем выполнения обратной функциональной операции над принятым сигналом.

Нас здесь будет интересовать случай, когда сигнал при передаче испытывает не всегда одинаковые изменения. В этом случае можно считать, что принятый сигнал E является функцией переданного сигнала S и второй переменной — шума N :

$$E = f(S, N).$$

Шум рассматривается как случайная переменная, точно так же, как выше рассматривалось сообщение. В общем случае шум может быть представлен соответствующим стохастическим процессом. Наиболее общий тип дискретного канала с шумом, который мы рассмотрим, является обобщением ранее описанного канала без шума с конечным числом состояний. Предположим, что число состояний конечно, и имеется множество вероятностей:

$$P_{\alpha, i}(\beta, j).$$

Это есть вероятность того, что если канал находится в состоянии α и передается символ i , то будет принят символ j и канал перейдет в состояние β . Таким образом, α и β пробегает все возможные состояния, i — все возможные передаваемые сигналы, а j — все возможные принимаемые сигналы. В том случае, когда последовательно передаваемые символы искажаются шумом независимо, имеется только одно состояние и канал описывается множеством переходных вероятностей $p_i(j)$ (вероятностей того, что переданный символ i будет принят как j).

Если канал с шумом питается некоторым источником, то имеются два статистических процесса: источник и шум. Поэтому имеется несколько энтропий, которые могут быть вычислены. Во-первых, существует энтропия источника или энтропия входа канала $H(x)$ (они равны, если передатчик невырожденный). Энтропия выхода канала, т. е. принятого сигнала, будет обозначаться через $H(y)$. В случае отсутствия шума $H(x) = H(y)$. Совместную энтропию входа и выхода обозначим $H(x, y)$. Наконец, имеются две условные энтропии $H_x(y)$ и $H_y(x)$ (энтропия выхода, когда вход известен, и наоборот). Эти величины связаны соотношениями

$$H(x, y) = H(x) + H_x(y) = H(y) + H_y(x).$$

Все эти энтропии могут измеряться как энтропии на одну секунду или на один символ.

12. Надежность и пропускная способность канала

При наличии шума, вообще говоря, невозможно восстановить с полной определенностью исходное сообщение или переданный сигнал, применяя какую-нибудь операцию к принятому сигналу E . Имеются, однако, некоторые способы передачи информации, которые являются оптимальными в отношении борьбы с шумом. Этот вопрос мы и рассмотрим теперь.

Предположим, что имеются два возможных символа 0 и 1, и мы ведем передачу со скоростью 1000 символов в секунду с вероятностями $p_0 = p_1 = \frac{1}{2}$. Таким образом, наш источник создает

информацию со скоростью 1000 бит в секунду. В процессе передачи шум вносит ошибки, так что в среднем один из ста символов принимается неправильно (0 вместо 1 или 1 вместо 0). Какова скорость передачи информации? Ясно, что она меньше 1000 бит в секунду, так как около 1% принятых сигналов неправильны. Первое желание состоит в том, чтобы сказать, что эта скорость составляет 990 бит в секунду, т. е. просто вычесть среднее число ошибок. Однако это не совсем правильно, так как при этом не учитывается, что получатель не знает, где именно произошла ошибка. Можно рассмотреть крайний случай и предположить, что шум настолько велик, что принятые символы полностью не зависят от переданных. Вероятность принять 1 (0) равна $1/2$, какой бы символ не был передан (1 или 0). Тогда благодаря одной случайности около половины принятых символов будут правильными, и нам пришлось бы считать, что система способна передавать 500 бит в секунду. Однако в действительности никакой передачи информации нет вовсе. Можно было бы получить столь же «хорошую» передачу, обойдясь вообще без канала и подбрасывая монету в точке приема.

Очевидно, истинная поправка к количеству переданной информации равна той части этой информации, которая отсутствует в принятом сигнале, или иначе той неопределенности относительно переданного сигнала, которая имеет место, когда известен принятый сигнал. Исходя из проведенного нами обсуждения понятия энтропии как меры неопределенности, представляется рациональным использовать условную энтропию сообщения (при известном сигнале) как меру этой недостающей части информации. Как будет видно ниже, такое определение действительно является подходящим. Следуя этой идее, было бы можно получить скорость действительной передачи информации R , вычитая из скорости создания информации (т. е. энтропии источника) среднюю скорость условной энтропии

$$R = H(x) - H_y(x).$$

Условная энтропия $H_y(x)$ для удобства будет называться *ненадежностью*¹⁾. Она является мерой средней неопределенности принятого сигнала.

В рассмотренном выше примере, если принят символ 0, то апостериорная вероятность того, что был передан символ 0, равна 0,99, а того, что была передана единица — 0,01. Если же была принята единица, то эти цифры поменяются местами. Следовательно,

$$H_y(x) = -[0,99 \log 0,99 + 0,01 \log 0,01] = 0,081 \text{ бит/символ}$$

¹⁾ Соответствующий английский термин «equivocation» переводится иногда в русской литературе как «неопределенность». — Прим. ред.

или 81 бит в секунду. Мы можем сказать, что система передает информацию со скоростью $1000 - 81 = 919$ бит в секунду. В том крайнем случае, когда при передаче символа 0 равновероятен прием как 0, так и 1, и аналогично для символа 1, апостериорные вероятности равны $1/2$, $1/2$ и тогда:

$$H_y(x) = - \left[\frac{1}{2} \log \frac{1}{2} + \frac{1}{2} \log \frac{1}{2} \right] = 1 \text{ бит/символ}$$

или 1000 бит в секунду. В этом случае скорость передачи равна 0, как и следовало ожидать.

Следующая теорема дает непосредственную интуитивную интерпретацию ненадежности, а также служит тому, чтобы оправдать ее как единственную подходящую меру. Рассмотрим систему связи и наблюдателя (или вспомогательный прибор), который может видеть как то, что передается, так и то, что принимается (с ошибками из-за шума). Этот наблюдатель отмечает ошибки в принятом сообщении и передает эти данные в точку приема через «коррекционный канал», чтобы дать возможность в точке приема исправить ошибки. Схематически это показано на рис. 8.

Теорема 10. Если коррекционный канал имеет пропускную способность, равную $H_y(x)$, то можно так закодировать данные коррекции, чтобы их можно было передавать по этому каналу и корректировать все ошибки, за исключением произвольно малой доли ϵ этих ошибок. Это невозможно, если пропускная способность коррекционного канала меньше чем $H_y(x)$.

Грубо говоря, $H_y(x)$ есть количество дополнительной информации, которая должна передаваться в точку приема за одну секунду для исправления принятого сообщения.

Для доказательства первой части рассмотрим длинные последовательности принятых сообщений M' и соответствующие исходные сообщения M . Для каждого из сообщений M' среди сообщений M будут иметься (с логарифмической точностью) $TH_y(x)$ сообщений, из которых оно могло бы быть по существу создано. Поэтому требуется передавать каждые T секунд $TH_y(x)$ бит. Это может быть осуществлено с частотой ошибок ϵ при помощи канала с пропускной способностью $H_y(x)$.

Вторая часть теоремы может быть доказана, если заметить, что для любых дискретных случайных переменных x, y, z

$$H_y(x, z) \geq H_y(x).$$

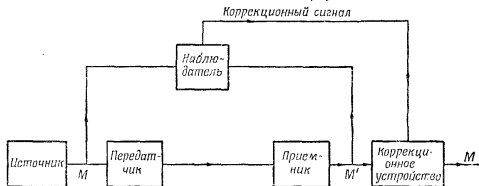
Разлагая дальше левую часть, получим

$$H_y(z) + H_{yz}(x) \geq H_y(x),$$

$$H_{yz}(x) \geq H_y(x) - H_y(z) \geq H_y(x) - H(z).$$

Если отождествить x с выходным сигналом источника, y с принятым сигналом и z с сигналом, посланным по коррекционному каналу, то правая часть равна ненадежности за вычетом скорости передачи по коррекционному каналу. Если пропускная способность этого канала меньше ненадежности, то правая часть будет больше нуля и $H_{yz}(x) > 0$. Но это выражение есть неопределенность того, что было передано при известном принятом сигнале и корректирующем сигнале. Если эта неопределенность больше нуля, то частота ошибок не может быть сделана произвольно малой.

Пример. Предположим, что в двоичной последовательности ошибки происходят случайным образом: вероятность того, что символ неправильный, равна p и вероятность того, что символ



Р и с. 8. Схема корректирующей системы.

правильный, $q = 1 - p$. Эти ошибки могут быть исправлены, если известны их положения. Таким образом, по коррекционному каналу нужно посылать информацию только об этих позициях. Это сводится к передаче информации от источника, который создает двоичные цифры, причем 1 имеет вероятность p (неправильно) и 0 — вероятность q (правильно). Для этого требуется канал с пропускной способностью

$$-[p \log p + q \log q],$$

что равно ненадежности исходной системы.

Исходя из упомянутых выше тождеств, скорость передачи R может быть записана в двух других формах. А именно:

$$R = H(x) - H_y(x) = H(y) - H_x(y) = H(x) + H(y) - H(x, y).$$

Первое из этих выражений уже было интерпретировано как количество посылаемой информации минус неопределенность того, что было послано. Второе выражение является мерой количества принятой информации минус та ее часть, которая обусловлена шумом. Третье выражение есть сумма этих двух количеств минус совместная энтропия, и поэтому оно в некотором смысле предстает

вляет число бит в секунду, общее этим двум количествам. Таким образом, все три выражения имеют определенный интуитивный смысл.

Пропускная способность канала с шумом должна быть максимально возможной скоростью передачи, т. е. скоростью при надлежащем согласовании источника с каналом. Определим поэтому пропускную способность канала как $C = \max [H(x) - H_y(x)]$, где максимум берется по всем возможным источникам информации, используемым в качестве входа в канал. Если рассматривается канал без шума, то $H_y(x) = 0$. Тогда это определение эквивалентно уже данному ранее определению для канала без шума, так как по теореме 8 максимум энтропии для канала равен его пропускной способности.

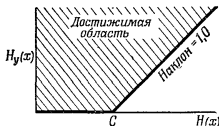
13. Основная теорема для дискретного канала с шумом

Может показаться неожиданным введение определенной пропускной способности C для канала с шумом, ибо в этом случае невозможно передавать информацию с достоверностью. Однако ясно, что, посылая информацию в избыточной форме, можно уменьшить вероятность ошибок. Например, путем многократного повторения сообщения и путем статистического изучения различных принятых вариантов сообщения вероятность ошибок может быть сделана очень малой. Можно было бы ожидать, однако, что, для того чтобы приблизить вероятности ошибок к нулю, нужно неограниченно увеличивать избыточность кодирования, и, следовательно, скорость передачи будет стремиться к нулю. Это, однако, ни в коей мере не верно. Если бы это было так, то не существовало бы вполне определенной пропускной способности, а была бы пропускная способность при заданной частоте ошибок или при заданной ненадежности, причем пропускная способность уменьшалась бы по мере того, как требования относительно ошибок становились бы все более жесткими. В действительности определенная выше пропускная способность C имеет вполне определенное значение. При надлежащем кодировании по каналу можно передавать информацию со скоростью C со сколь угодно малой частотой ошибок или при сколь угодно малой ненадежности. Это утверждение неверно, если скорость передачи больше C . При попытках передавать со скоростью больше чем C , скажем $C + R_1$, неизбежно появится ненадежность, равная или большая чем R_1 . Природа берет свою плату, вводя столь большую неопределенность, что в действительности при передаче скоростью большей C , нельзя получить полностью не искаженную информацию.

Изложенное выше показано на рис. 9. Скорость информации в канале отложена по горизонтали, а ненадежность — по вертикали.

Любая точка выше жирной линии в заштрихованной области может быть получена, точки же ниже линии получены быть не могут. Точки самой линии, вообще говоря, получены быть не могут, за исключением обычно двух точек.

Эти положения являются основным оправданием предложенного определения C и будут сейчас доказаны.



Р и с. 9. Ненадежность, возможная для энтропии входа канала.

Т е о р е м а 11. Пусть дискретный канал обладает пропускной способностью C , а дискретный источник — энтропией в секунду H . Если $H \leq C$, то существует такая система кодирования, что сообщения источника могут быть переданы по каналу с произвольно малой частотой ошибок (или со сколь угодно малой ненадежностью). Если $H > C$, то можно закодировать источник таким образом, что ненадежность будет меньше чем $H - C + \epsilon$, где ϵ сколь угодно мало. Не существует способа кодирования, обеспечивающего ненадежность, меньшую чем $H - C^1$.

Метод доказательства первой части этой теоремы состоит не в указании способа кодирования, имеющего желаемые свойства, а в доказательстве того, что искомый код должен существовать в определенной группе кодов. Будем усреднять по этой группе частоту ошибок и покажем, что полученное среднее может быть сделано меньше чем ϵ . Но если среднее некоторого множества чисел меньше чем ϵ , то в этом множестве должно существовать по крайней мере одно число, меньшее ϵ . Это и даст желаемый результат.

Пропускная способность канала с шумом была определена как

$$C = \max [H(x) - H_y(x)],$$

где x — есть сигнал на входе канала, а y — на выходе канала. Максимум берется по всем источникам, которые могут быть использованы на входе такого канала.

¹) Математически полное доказательство этой теоремы иным методом можно найти в книге: Ф а й н с т е й н А., Теория информации, ИЛ, М., 1960, а в более общей постановке и с использованием метода, развитого здесь автором, в статье: Д о б р у ш и н Р. Л., Общая формулировка основной теоремы Шеннона в теории информации, Успехи матем. наук (1959), 14, 6, 3. — Прим. ред.

Пусть S_0 — источник, на котором достигается максимальная пропускная способность C . Если максимум в действительности не достигается ни на каком источнике (но достигается лишь в пределе), то пусть S_0 — источник, для которого пропускная способность достаточно близка к предельной.

Пусть S_0 используется в качестве входа канала. Рассмотрим возможные передаваемые и принимаемые последовательности большой длительности T . Можно утверждать следующее:

1. Передаваемые последовательности распадаются на два класса: высоковероятная группа, содержащая примерно $2^{TH(x)}$ членов, и остающиеся последовательности с малой суммарной вероятностью.

2. Аналогично имеется высоковероятное множество принимаемых последовательностей, содержащее примерно $2^{TH(y)}$ членов, и маловероятное множество оставшихся последовательностей.

3. Каждая из высоковероятных выходных последовательностей может быть создана одной из входных последовательностей, число которых равно примерно $2^{TH_y(x)}$. Суммарная вероятность всех других случаев мала.

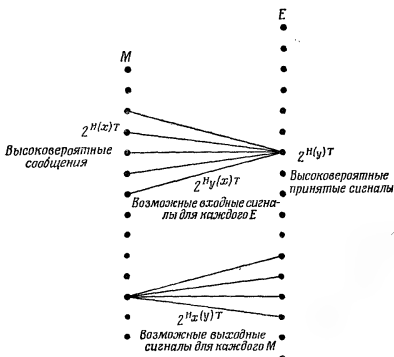
4. Каждая из высоковероятных входных последовательностей может создать примерно $2^{TH_x(y)}$ выходных последовательностей. Суммарная вероятность всех других исходов мала.

В этих утверждениях все «ε» и «δ», подразумеваемые в словах «малый» и «примерно», стремятся к нулю, когда T возрастает и S_0 приближается к максимизирующему источнику.

Вышесказанное изображено на рис. 10, где входные последовательности являются точками слева, а выходные последовательности — точками справа. Верхний веер сходящихся линий изображает ряд возможных входов для типичного выхода. Нижний веер изображает ряд возможных результатов, обусловленных типичным входом. В обоих случаях не учитываются маловероятные множества.

Предположим теперь, что имеется другой источник S , создающий информацию со скоростью R , причем $R < C$. За время T этот источник будет создавать 2^{TR} высоковероятных сообщений. Требуется связать их с некоторым выбранным множеством возможных входных сигналов канала таким образом, чтобы получить малую частоту ошибок. Будем устанавливать эту связь всеми возможными способами (используя, однако, только ту высоковероятную группу входных сигналов, которая определяется источником S_0) и усредним частоту ошибок по этому большому классу возможных систем кодирования. Это равносильно вычислению частоты ошибок при случайном связывании сообщений и входных сигналов канала длительности T . Предположим, что наблюдается некоторый выходной сигнал y_1 . Какова вероятность того, что во множество возможных причин получения этого y_1 войдет более

одного сообщения от источника S . Имеется 2^{TR} сообщений, распределенных случайно по $2^{TH(x)}$ точкам. Вероятность того, что



Р и с. 10. Схематическое представление соотношений между входами и выходами в канале.

некоторая конкретная точка будет сообщением, поэтому равна $2^{T[R-H(x)]}$.

Вероятность того, что никакая другая точка всеера (кроме действительного исходного сообщения) не будет сообщением, равна

$$P = [1 - 2^{T[R-H(x)]}] 2^{TH_y(x)}.$$

Но $R < H(x) - H_y(x)$, так что $R - H(x) = -H_y(x) - \eta$,

где η положительно. Следовательно,

$$P = [1 - 2^{-TH_y(x) - T\eta}] 2^{TH_y(x)}$$

стремится (при $T \rightarrow \infty$) к

$$1 - 2^{-T\eta}.$$

Отсюда вероятность ошибки стремится к нулю, и первая часть теоремы доказана.

Вторую часть теоремы легко доказать исходя из того, что можно просто передавать C бит в секунду от источника, совсем пренебрегая остатком создаваемой информации. В приемнике эта неучитываемая часть создаст ненадежность $H(x) - C$, а передаваемая часть должна добавить лишь ε . Эту границу можно достигнуть также многими другими способами, как будет доказано при рассмотрении случая непрерывного канала.

Последнее утверждение теоремы является простым следствием нашего определения C . Предположим, что можно закодировать некоторый источник с энтропией $H(x) = C + a$ таким образом, чтобы получить ненадежность $H_y(x) = a - \varepsilon$, где ε положительно. Тогда

$$H(x) - H_y(x) = C + \varepsilon,$$

причем ε положительно. Это противоречит определению C как максимума $H(x) - H_y(x)$.

В действительности здесь доказано больше, чем утверждалось в теореме. Если среднее значение множества положительных чисел заключено между нулем и ε , то только часть из них, доля которой не превышает $\sqrt{\varepsilon}$, может быть больше $\sqrt{\varepsilon}$. Так как ε произвольно мало, то можно сказать, что почти все системы кодирования сколь угодно близки к идеальной.

14. Рассмотрение полученных результатов

Доказательство теоремы 11, не будучи чистым доказательством существования, обладает некоторыми недостатками подобных доказательств. Попытка осуществить хорошее приближение к идеальному кодированию по методу, примененному в доказательстве, вообще говоря, представляется непрактичной. Действительно, за исключением нескольких довольно тривиальных случаев и некоторых предельных ситуаций, никакого явного описания последовательных приближений к идеальному методу не найдено. Вероятно, это не случайно, а связано с трудностью задания явного построения хорошей аппроксимации случайной последовательности.

Приближение к идеальному методу передачи обладало бы тем свойством, что если сигнал изменен шумом в умеренных пределах, то оригинал все еще может быть восстановлен. Другими словами, изменение не делает, вообще говоря, принимаемый сигнал ближе к другим возможным сигналам, чем к оригиналу. Это достигается ценой введения определенной избыточности в кодировании. Избыточность должна быть введена способом, приспособленным для борьбы против действующих на канал шумов определенной структуры. Впрочем, обычно будет помогать любая избыточность источника, если она используется в точке приема. В частности, если

источник уже имеет некоторую избыточность и не делается никаких попыток устранить ее при согласовании с каналом, эта избыточность будет помогать борьбе с шумом. Например, в телеграфном канале без шума можно сэкономить около 50% времени с помощью надлежащего кодирования сообщений. Этого не делается, и большая часть избыточности английского текста остается и в символах канала. Впрочем, это имеет и преимущество, так как оказывается допустимым значительный шум в канале. Значительная часть букв может приниматься неправильно и все же восстанавливаться на основании контекста. В действительности, при этом во многих случаях получается, по-видимому, неплохое приближение к идеальному кодированию, так как статистическая структура английского текста является довольно запутанной, и осмысленные английские последовательности не слишком далеки (в смысле, требуемом для теоремы) от случайного выбора.

Как и в случае отсутствия шума, для приближения к идеальному кодированию требуется, вообще говоря, некоторая временная задержка. Теперь она приобретает новую функцию, позволяя воздействовать на сигнал большой выборкой шума до того, как будет сделано какое-либо суждение в точке приема относительно исходного сообщения. Увеличение объема выборки всегда уточняет возможные статистические утверждения.

Содержание теоремы 11 и ее доказательство могут быть сформулированы несколько иным способом, который яснее выявляет связь со случаем отсутствия шума. Рассмотрим возможные сигналы длительности T и предположим, что из них выбирается для использования некоторое подмножество. Пусть все сигналы этого подмножества используются с одинаковой вероятностью, и предположим, что приемник сконструирован так, что он выбирает в качестве исходного сигнала тот элемент из нашего подмножества, для которого наиболее вероятно перейти в искаженный сигнал. Обозначим через $N(T, q)$ максимальное число сигналов, которые можно выбрать для нашего подмножества так, чтобы вероятность неправильной интерпретации была меньше или равна q .

Теорема 12. Если q не равно 0 или 1, то

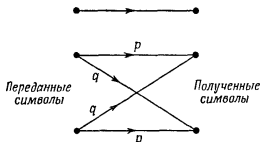
$$\lim_{T \rightarrow \infty} \frac{\log N(T, q)}{T} = C,$$

где C — пропускная способность канала.

Другими словами, независимо от требований надежности можно за время T уверенно различить достаточное количество сообщений, соответствующее примерно CT битам, если T достаточно велико. Теорему 12 можно сравнить с определением пропускной способности канала без шума, данным в первой части статьи.

15. Пример дискретного канала и его пропускной способности

Простой пример дискретного канала показан на рис. 11. Имеются три возможных символа. Первый символ не подвергается воздействию шума. Второй и третий символы имеют вероятность p пройти неискаженными и вероятность q превратиться в другой символ



Р и с. 11. Пример дискретного канала.

той же пары. Пусть $\alpha = -[p \log p + q \log q]$ и пусть P , Q и Q — вероятности передачи соответственно первого, второго и третьего символов (при этом две последние вероятности равны по соображениям симметрии). Имеем

$$H(x) = -P \log P - 2Q \log Q,$$

$$H_y(x) = 2Q\alpha.$$

Требуется выбрать P и Q так, чтобы максимизировать $H(x) - H_y(x)$, соблюдая при этом условие $P + 2Q = 1$. Поэтому рассмотрим

$$U = -P \log P - 2Q \log Q - 2Q\alpha + \lambda (P + 2Q),$$

$$\frac{\partial U}{\partial P} = -1 - \log P + \lambda = 0,$$

$$\frac{\partial U}{\partial Q} = -2 - 2 \log Q - 2\alpha + 2\lambda = 0.$$

Исключая λ , получим

$$\log P = \log Q + \alpha,$$

$$P = Q \cdot 2^\alpha = Q\beta,$$

$$P = \frac{\beta}{\beta+2}, \quad Q = \frac{1}{\beta+2}.$$

Пропускная способность канала равна

$$C = \log \frac{\beta+2}{\beta}.$$

Заметим, что полученные выражения дают очевидные ответы в случаях $p = 1$ и $p = \frac{1}{2}$. В первом случае $\beta = 1$ и $C = \log 3$,

что правильно, так как в этом случае имеется канал без шума с тремя возможными символами. Если $p = \frac{1}{2}$, то $\beta = 2$ и $C = \log 2$. Здесь второй и третий символы не могут быть отличены друг от друга и они воспринимаются как один символ. Первый символ передается с вероятностью $P = \frac{1}{2}$, и второй с третьим вместе — с вероятностью $\frac{1}{2}$. Эта вероятность может быть распределена между вторым и третьим символами произвольным образом, и при этом всегда будет достигаться максимальная пропускная способность.

При промежуточных значениях p пропускная способность канала будет заключена между $\log 2$ и $\log 3$. Различие между вторым и третьим символами несет некоторое количество информации, но не так много, как в случае отсутствия шума. Первый символ передается несколько чаще, чем два остальных, так как на него не воздействует шум.

16. Пропускная способность канала в некоторых частных случаях

Если шум воздействует на последовательные символы в канале независимо, то канал может быть описан множеством переходных вероятностей p_{ij} , где p_{ij} есть вероятность того, что если послан символ i , то будет принят символ j . Пропускная способность канала равна в этом случае максимуму выражения

$$-\sum_{i,j} P_i p_{ij} \log \sum_i P_i p_{ij} + \sum_{i,j} P_i p_{ij} \log p_{ij},$$

где мы варьируем P_i , соблюдая условие $\sum_i P_i = 1$. Применение метода Лагранжа приводит нас к уравнениям

$$\sum_j p_{sj} \log \frac{p_{sj}}{\sum_i P_i p_{ij}} = \mu \quad s = 1, 2, \dots$$

Умножение на P_s и суммирование по s показывают, что $\mu = -C$. Обозначим обратную матрицу к p_{sj} (если она существует) через h_{st} , так что $\sum_s h_{st} p_{sj} = \delta_{it}$. Тогда

$$\sum_{s,j} h_{st} p_{sj} \log p_{sj} - \log \sum_i P_i p_{it} = -C \sum_s h_{st}.$$

Отсюда

$$\sum_i P_i p_{it} = \exp \left\{ C \sum_s h_{st} + \sum_{s,j} h_{st} p_{sj} \log p_{sj} \right\}$$

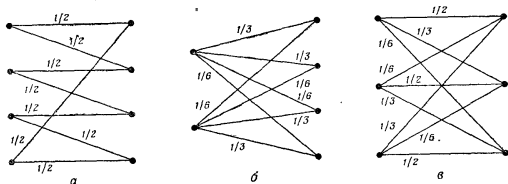
или

$$P_i = \sum_t h_{it} \exp \left\{ C \sum_s h_{st} + \sum_{s,j} h_{st} p_{sj} \log p_{sj} \right\}.$$

Это — система уравнений для определения максимизирующих значений P_i , причем C должно быть определено таким образом, чтобы $\sum_i P_i = 1$. Когда это выполнено, C будет пропускной способ-

ностью канала, а P_i — соответствующими вероятностями для символов канала, для которых достигается эта пропускная способность¹⁾.

Если каждый входной символ имеет одинаковое множество вероятностей, изображенных на схеме исходящими от него линиями, и это же справедливо для каждого выходного символа, пропускная



Р и с. 12. Примеры дискретных каналов с теми же самыми переходными вероятностями для каждого входа и каждого выхода.

способность может быть легко вычислена. Примеры показаны на рис. 12. В таком случае $H_x(y)$ не зависит от распределения вероятностей между входными символами и равняется — $\sum p_i \log p_i$, где p_i — значения переходных вероятностей от любого входного символа. Пропускная способность канала равна

$$\text{Max} [H(y) - H_x(y)] = \text{Max} H(y) + \sum p_i \log p_i.$$

Максимум $H(y)$, очевидно, равен $\log m$, где m — число выходных символов, так как все они могут быть сделаны равновероятными, если сделать равновероятными входные символы. Поэтому пропускная способность канала равна

$$C = \log m + \sum p_i \log p_i.$$

На рис. 12, а пропускная способность C была бы равна

$$C = \log 4 - \log 2 = \log 2.$$

¹⁾ Дополнительная трудность состоит в том, что некоторые из P_i , вычисленные по этой формуле, могут оказаться отрицательными. Эти вопросы подробно разобраны в статье Мурога [Muroga S., On the capacity of a discrete channel, *J. Phys. Soc. Japan* (1956)]. — Прим. ред.

Это значение могло бы быть достигнуто при использовании лишь первого и третьего символов. На рис. 12, б имеем:

$$\begin{aligned} C &= \log 4 - \frac{2}{3} \log 3 - \frac{1}{3} \log 6 = \\ &= \log 4 - \log 3 - \frac{1}{3} \log 2 = \\ &= \log \frac{2^{5/3}}{3}. \end{aligned}$$

На рис. 12, в имеем:

$$C = \log 3 - \frac{1}{2} \log 2 - \frac{1}{3} \log 3 - \frac{1}{6} \log 6 = \log \frac{3}{2^{1/2} \cdot 3^{1/3} \cdot 6^{1/6}}.$$

Предположим, что символы распадаются на различные группы таким образом, что шум не превращает символы одной группы в символы другой группы. Пусть пропускная способность для n -й группы равна C_n (в битах в секунду), если передаются только символы этой группы. Тогда легко показать, что для наилучшего использования всей совокупности суммарная вероятность p_n всех символов в n -й группе должна быть равна

$$P_n = \frac{2^{C_n}}{\sum_n 2^{C_n}}.$$

Внутри группы эти вероятности распределяются как раз в такой пропорции, как если бы они были единственно используемыми. Пропускная способность канала равна

$$C = \log \sum 2^{C_n}.$$

17. Пример эффективного кодирования

Следующий пример, хотя и несколько искусственный, представляет случай, когда возможно точное согласование передатчика с каналом. В канале имеются два символа 0 и 1, а шум воздействует на блоки из семи символов. Блок из семи символов либо передается без ошибок, либо в нем оказывается ошибочным ровно один символ из семи. Все эти восемь возможностей равновероятны.

Имеем

$$C = \text{Max} [H(y) - H_x(y)] = \frac{1}{7} \left[7 + \frac{8}{8} \log \frac{1}{8} \right] = \frac{4}{7} \text{ бит/символ.}$$

Эффективный код, обеспечивающий полную коррекцию ошибок и передачу со скоростью C , представляет собой следующее (он найден по методу, предложенному Р. Хэммингом).

Пусть блок из семи символов будет $X_1 X_2, \dots, X_7$. Из них X_3, X_5, X_6, X_7 — символы сообщения и выбираются произвольно источником. Остальные три символа являются избыточными и вычисляются следующим образом:

X_4 выбирается так, чтобы $\alpha = X_4 + X_5 + X_6 + X_7$,

X_2 ————— » ————— » ————— $\beta = X_2 + X_3 + X_6 + X_7$,

X_1 ————— » ————— » ————— $\gamma = X_1 + X_3 + X_5 + X_7$ были четными.

Когда принят блок из семи символов, вычисляются α, β, γ и если какое-либо из них окажется четным, то считаем его нулем, если же нечетным, то единицей. Двоичное число $\alpha\beta\gamma$ даст тогда индекс того X_i , которое оказалось ошибочным (если получится 0, то блок принят без ошибок)¹⁾.

III. ИНФОРМАЦИЯ ДЛЯ НЕПРЕРЫВНЫХ ВЕЛИЧИН

Рассмотрим теперь случай, когда сигналы или сообщения (или те и другие) являются непрерывными величинами в отличие от предполагавшейся до этого дискретности рассматриваемых систем. В значительной степени результаты для непрерывного случая могут быть получены предельным переходом от дискретного случая путем деления всего континуума сообщений или сигналов на большое, но конечное число малых областей и вычисления различных параметров, введенных на дискретной основе. По мере уменьшения размеров областей эти параметры, вообще говоря, сходятся в пределе к соответствующим значениям для непрерывного случая. Однако появляется несколько новых эффектов, а также наблюдается общее перемещение центра тяжести в сторону перехода от общих результатов к специальным частным случаям.

В непрерывных системах не будет делаться попыток получать результаты с наибольшей общностью или с крайней строгостью чистой математики, так как это потребовало бы серьезного применения абстрактной теории меры и затемнило бы основную нить нашего анализа. Предварительное изучение, однако, показывает, что теория может быть сформулирована полностью аксиоматическим и строгим образом, включая в себя непрерывный дискретный и многие другие случаи. Некоторые вольности, допущенные в настоящем анализе, при предельных переходах могут быть оправданы во всех случаях, представляющих практический интерес.

¹⁾ Некоторые дальнейшие примеры самокорректирующихся кодов см. в работе Goley M. J. F., Notes on digital coding, *Proc. of the Institute of Radio Engineers*, 37, № 6, June (1949), 637.

18. Множества и ансамбли функций

В непрерывных системах мы будем иметь дело с множествами функций и ансамблями функций. Множество функций, как указывает сам термин, есть просто некоторый класс или набор функций обычно от одной переменной — времени. Оно может быть определено либо путем явного представления различных функций множества, либо неявно путем указания тех свойств, которыми функции множества обладают, а другие функции не обладают. Приведем некоторые примеры:

1. Множество функций

$$f_{\theta}(t) = \sin(t + \theta).$$

Каждое частное значение θ определяет частную функцию множества.

2. Множество всех функций времени, не содержащих частот выше W герц.

3. Множество всех функций, ограниченных по полосе числом W и по амплитуде числом A .

4. Множество всех английских речевых сигналов, рассматриваемых как функции времени.

Ансамбль функций¹⁾ есть множество функций вместе с вероятностной мерой, посредством которой можно определить вероятность того, что функция множества обладает определенными свойствами²⁾. Например, вместе с множеством

$$f_{\theta}(t) = \sin(t + \theta)$$

можно задать распределение вероятностей для θ , скажем $P(\theta)$. Тогда это множество становится ансамблем.

Приведем некоторые другие примеры ансамблей функций:

1. Конечное множество функций $f_k(t)$, ($k = 1, 2, \dots, n$), где f_k имеет вероятность p_k .

2. Конечнопараметрическое семейство функций

$$f(a_1, a_2, \dots, a_n; t)$$

с распределением вероятностей для параметров a_i

$$p(a_1, \dots, a_n).$$

Например, ансамбль, определяемый выражением

$$f(a_1, \dots, a_n, \theta_1, \dots, \theta_n; t) = \sum_{k=1}^n a_n \sin k(\omega t + \theta_n),$$

¹⁾ Ансамбль функций называют обычно вероятностным процессом. Доступное изложение понятий теории вероятностных процессов можно, например, найти в книге Давенпорта и Рута (Д а в е н п о р т В. Б., Р у т В. Л., Введение в теорию случайных сигналов и шумов, ИЛ, М., 1960).

²⁾ Точнее говоря, эти функции принадлежат к измеримому пространству с полной мерой пространства, равной единице.

причем амплитуды a_i распределены нормально и независимы, а фазы θ_i распределены равномерно в интервале $(0, 2\pi)$ и независимы.

3. Ансамбль

$$f(a_i, t) = \sum_{n=-\infty}^{\infty} a_n \frac{\sin \pi (2Wt - n)}{\pi (2Wt - n)},$$

где a_i распределены нормально и независимо и все имеют одно и то же стандартное отклонение $\sqrt{N}$. Это есть одно из представлений «белого» шума с полосой частот от 0 до W герц и со средней мощностью N^1).

4. Пусть на оси t распределены точки по закону Пуассона. В каждую выбранную точку помещается функция $f(t)$, и различные функции складываются, давая ансамбль

$$\sum_{k=-\infty}^{\infty} f(t + t_k),$$

где t_k — точки, распределенные по закону Пуассона. Этот ансамбль может рассматриваться как разновидность импульсных или дробовых шумов, когда все импульсы одинаковы.

5. Система английских речевых функций с вероятностной мерой, определяемой частотой их появления при их повседневном использовании.

Ансамбль функций $f_\theta(t)$ называется *стационарным*, если при сдвиге всех функций по времени на некоторую фиксированную величину получается тот же самый ансамбль. Ансамбль

$$f_\theta(t) = \sin(t + \theta)$$

является стационарным, если θ распределено равномерно от 0 до 2π . Если сдвинуть каждую функцию на t_1 , то получится

$$f_\theta(t + t_1) = \sin(t + t_1 + \theta) = \sin(t + \varphi),$$

где φ распределено равномерно от 0 до 2π . Каждая функция изменилась, но ансамбль в целом при этом сдвиге остался неизменным. В других примерах, приведенных выше, ансамбли также стационарны.

¹) Это представление может быть использовано как определение белого шума с ограниченной полосой частот. Оно имеет некоторые преимущества, так как содержит меньшее число предельных переходов, нежели определения, применявшиеся ранее. Термин «белый шум», уже прочно укоренившийся в литературе, представляется несколько неудачным. В оптике под белым светом понимается или свет с каким-либо непрерывным спектром (в противоположность линейчатому), или же со спектром, равномерным по отношению к длине волны (а это не то же самое, что спектр, равномерный по отношению к частоте).

Ансамбль называется эргодическим, если он является стационарным и если во множестве функций не существует стационарного подмножества с вероятностью, отличной от 0 и 1. Ансамбль

$$\sin(t + \theta)$$

является эргодическим. Никакое подмножество этих функций с вероятностью, отличной от 0 и 1, не может быть превращено само в себя при всех временных сдвигах. С другой стороны, ансамбль

$$a \sin(t + \theta),$$

где a распределено нормально, а θ — равномерно, является стационарным, но не эргодическим. Подмножество таких функций, для которого a заключено, например, между 0 и 1, является стационарным и имеет вероятность, не равную 0 и 1.

Из приведенных выше примеров ансамблей 3-й и 4-й являются эргодическими, а 5-й, возможно, также может рассматриваться как эргодический. Если ансамбль эргодический, то, грубо говоря, каждая функция множества является типичной для ансамбля. Более точно известно, что для эргодического ансамбля среднее любой статистики по ансамблю равно (с вероятностью, равной единице) ее среднему по всем временным переносам любой частной функции множества¹⁾. Грубо говоря, можно ожидать, что при изменении времени каждая функция испытает с надлежащей частотой все изменения, претерпеваемые любой из функций множества.

Выполняя различные операции над числами или функциями, можно получить новые числа или функции. Точно так же можно выполнять операции над ансамблями для получения новых ансамблей. Предположим, например, что имеется ансамбль функций $f_\alpha(t)$ и оператор T , переводящий каждую функцию $f_\alpha(t)$ в $g_\alpha(t)$:

$$g_\alpha(t) = T f_\alpha(t).$$

Вероятностная мера для множества $g_\alpha(t)$ определяется вероятностной мерой для множества $f_\alpha(t)$. Вероятность некоторого подмножества функций $g_\alpha(t)$ равна вероятности такого подмножества функций $f_\alpha(t)$, члены которого переводятся оператором T в члены дан-

¹⁾ Это несть знаменитая эргодическая теорема или, скорее, один из аспектов этой теоремы, которая была доказана в несколько отличных формулировках Биркгофом, Нейманом и Купманом и впоследствии обобщена Винером, Хопфом, Гуревичем и др. Литература по эргодической теории довольно обширна, и читатель отсылается к статьям этих авторов за точными и общими формулировками: например, H o p f E., *Ergodentheorie*, *Erg. der Math. und ihrer Grenzgebiete*, 5; On causality statistics and probability, *J. of Math. and Phys.*, 13, № 1, 1934; W i e n e r N., The ergodic theorem, *Duke Math. J.*, 5, 1939 [см. также Д у б Д. Л., Вероятностные процессы, М.—Л., 1956, гл. X.—Прим. ред.].

ного подмножества функций g . Физически это соответствует прохождению ансамбля через некоторое устройство, например фильтр, выпрямитель или модулятор. Функции на выходе устройства образуют ансамбль $g_\alpha(t)$.

Устройство или оператор T будет называться инвариантным, если сдвиг входа просто сдвигает выход, т. е. если из равенства

$$g_\alpha(t) = T f_\alpha(t)$$

следует, что

$$g_\alpha(t + t_1) = T f_\alpha(t + t_1)$$

для всех $f_\alpha(t)$ и всех t_1 . Легко показать (см. приложение 5), что если T — инвариантный оператор, а входной ансамбль стационарный, то выходной ансамбль также стационарный. Подобным же образом, если входной ансамбль эргодический, то выходной ансамбль также будет эргодическим.

Фильтр или выпрямитель инвариантны при всех временных переносах. Операция модуляции не является инвариантной, так как фаза несущей создает определенную временную структуру. Однако модуляция инвариантна при всех переносах, кратных периоду несущей.

Винер указал на тесную связь между инвариантностью физических устройств при временных переносах и теорией Фурье¹⁾. Он показал, что если устройство линейно и инвариантно, то анализ методом Фурье является удобным математическим аппаратом для решения задачи.

Ансамбль функций представляет собой подходящее математическое представление сообщений, создаваемых непрерывным источником (например, речью), причем как для сигналов, создаваемых передатчиком, так и для мешающих шумов. Теория связи имеет дело, как подчеркнул Винер, не с операциями над конкретными функциями, а с операциями над ансамблями функций. Система связи конструируется не для определенной речевой функции и тем более не для синусоидальной волны, а для ансамбля речевых функций.

¹⁾ Теория связи многими из своих принципов и теорий обязана Винеру. Его классическая работа «The interpolation, extrapolation and smoothing of stationary time series», Wiley, 1949, содержит первую четкую формулировку теории связи как статистической проблемы, изучение операций над временными рядами. Хотя эта работа главным образом касается линейного предсказания и проблемы фильтрации, представляется важным сослаться на нее в связи с настоящей статьей. Можно отослать здесь также к его известной работе «Кибернетика», рассматривающей общие проблемы связи и контроля. (Русский перевод: В и н е р Н., Кибернетика, Советское Радио, М., 1958.— Прим. ред.)

19. Ансамбли функций с ограниченной полосой частот

Если функция времени $f(t)$ ограничена по полосе частот от 0 до W герц, то она полностью определяется заданием ее ординат в дискретной последовательности точек, отстоящих друг от друга на $\frac{1}{2W}$ секунду при помощи способа, указываемого следующими рассуждениями:

Теорема 13¹⁾. Пусть $f(t)$ не содержит частот, превышающих W . Тогда

$$f(t) = \sum_{-\infty}^{\infty} X_n \frac{\sin \pi (2Wt - n)}{\pi (2Wt - n)},$$

где

$$X_n = f\left(\frac{n}{2W}\right).$$

В этом разложении $f(t)$ представляется как сумма ортогональных функций. Коэффициенты X_n при различных слагаемых можно рассматривать как координаты в бесконечномерном «функциональном пространстве». В этом пространстве каждая функция соответствует точно одной точке, и каждая точка — одной функции.

Можно считать, что функция фактически ограничена временным интервалом T , если все ординаты X_n вне этого интервала времени равны нулю. В этом случае только $2TW$ координат отличны от нуля. Таким образом, функции, ограниченные полосой частот W и длительностью T , соответствуют точкам в пространстве $2TW$ измерений.

Подмножество функций с полосой частот W и длительностью T соответствует области в этом пространстве. Например, функции, полная энергия которых меньше или равна E , соответствуют точкам $2TW$ -мерной сферы радиуса $r = \sqrt{2WE}$.

Ансамбль функций с ограниченной полосой частот и ограниченной длительностью будет представляться распределением вероятностей $p(x_1, \dots, x_n)$ в соответствующем n -мерном пространстве. Если ансамбль не ограничен по времени, мы можем считать, что $2TW$ координат в данном интервале T представляют фактически часть функции, соответствующую интервалу T , и распределение $p(x_1, \dots, x_n)$ дает статистическую структуру ансамбля для интервалов такой длительности.

¹⁾ Доказательство этой теоремы и дальнейшее обсуждение см. в работе автора «Связь при наличии шума», стр. 433 настоящего сборника. Теоретико-информационное содержание этой теоремы, давно известной в теории функций, было впервые указано В. А. Котельниковым. — Прим. ред.)

20. Энтропия непрерывного распределения

Энтропия дискретного множества вероятностей $p_1, \dots, p_n$ была определена как

$$H = - \sum p_i \log p_i.$$

Аналогичным образом определим энтропию непрерывного распределения с функцией плотности распределения $p(x)$ как

$$H = - \int_{-\infty}^{\infty} p(x) \log p(x) dx.$$

В случае n -мерного распределения $p(x_1, \dots, x_n)$ имеем

$$H = - \int \dots \int p(x_1, \dots, x_n) \log p(x_1, \dots, x_n) dx_1, \dots, dx_n.$$

Если имеются два аргумента x и y (которые сами могут быть многомерными), совместная и условная энтропии $p(x, y)$ даются выражениями

$$H(x, y) = - \iint p(x, y) \log p(x, y) dx dy$$

и

$$H_x(y) = - \iint p(x, y) \log \frac{p(x, y)}{p(x)} dx dy,$$

$$H_y(x) = - \iint p(x, y) \log \frac{p(x, y)}{p(y)} dx dy,$$

где

$$p(x) = \int p(x, y) dy,$$

$$p(y) = \int p(x, y) dx.$$

Энтропия непрерывных распределений обладает большинством свойств (но не всеми) энтропии для дискретных распределений. В частности:

1. Если x ограничено некоторой областью объема v в пространстве своих значений, то $H(x)$ будет максимальна и равна $\log v$, когда $p(x)$ равна константе $\frac{1}{v}$ в этой области.

2. Для двух переменных x, y имеем

$$H(x, y) \leq H(x) + H(y),$$

причем равенство имеет место тогда и только тогда, когда x и y независимы, т. е. $p(x, y) = p(x) \cdot p(y)$ (за исключением, может быть, множества точек, имеющего нулевую вероятность).

3. Рассмотрим обобщенную операцию усреднения следующего типа:

$$p'(y) = \int a(x, y) p(x) dx,$$

где

$$\int a(x, y) dx = \int a(x, y) dy = 1, \quad a(x, y) \geq 0.$$

Тогда энтропия усредненного распределения $p'(y)$ больше или равна энтропии первоначального распределения $p(x)$.

4. Имеем

$$H(x, y) = H(x) + H_x(y) = H(y) + H_y(x)$$

и

$$H_x(y) \leq H(y).$$

5. Пусть $p(x)$ — одномерное распределение. Распределение $p(x)$, дающее максимальную энтропию при условии, что стандартное отклонение x фиксировано и равно σ , является гауссовским. Чтобы показать это, надо максимизировать

$$H(x) = - \int p(x) \log p(x) dx$$

при ограничениях

$$\sigma^2 = \int x^2 p(x) dx \quad \text{и} \quad 1 = \int p(x) dx.$$

Это сводится, как известно из вариационного исчисления, к максимизации

$$\int [- p(x) \log p(x) + \lambda p(x) x^2 + \mu p(x)] dx.$$

Условием для этого будет

$$-1 - \log p(x) + \lambda x^2 + \mu = 0,$$

и, следовательно (подбирая постоянные так, чтобы удовлетворить приведенным выше ограничениям), получим

$$p(x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{x^2}{2\sigma^2}}.$$

Аналогичным образом в n -мерном случае предположим, что моменты второго порядка плотности $p(x_1, \dots, x_n)$ фиксированы и равны A_{ij} :

$$A_{ij} = \int \dots \int x_i x_j p(x_1, \dots, x_n) dx_1, \dots, dx_n.$$

Находим (при помощи аналогичных вычислений), что максимум энтропии имеет место, когда $p(x_1, \dots, x_n)$ является n -мерным нормальным распределением с моментами второго порядка A_{ij} .

6. Энтропия одномерного нормального распределения со стандартным отклонением σ равна

$$H(x) = \log_2 \sqrt{2\pi} \sigma.$$

Это вычисляется следующим образом:

$$p(x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{x^2}{2\sigma^2}},$$

$$-\log_2 p(x) = \log_2 \sqrt{2\pi}\sigma + \frac{x^2}{2\sigma^2} \log_2 e,$$

$$\begin{aligned} H(x) &= -\int p(x) \log p(x) dx = \\ &= \int p(x) \log \sqrt{2\pi}\sigma dx + \int p(x) \frac{x^2}{2\sigma^2} dx = \\ &= \log \sqrt{2\pi}\sigma + \frac{\sigma^2}{2\sigma^2} = \log \sqrt{2\pi}\sigma + \log \sqrt{e} = \log \sqrt{2\pi e}\sigma. \end{aligned}$$

Аналогичным образом n -мерное гауссовское распределение с квадратичной формой a_{ij} задается как

$$p(x_1, \dots, x_n) = \frac{|a_{ij}|^{\frac{1}{2}}}{(2\pi)^{\frac{n}{2}}} \exp \left\{ -\frac{1}{2} \sum a_{ij} X_i X_j \right\},$$

а энтропия равна

$$H = \log (2\pi e)^{\frac{n}{2}} |a_{ij}|^{-\frac{1}{2}},$$

где $|a_{ij}|$ — определитель матрицы с элементами a_{ij} .

7. Если распределение x ограничено положительной полуосью [$p(x) = 0$ для $x \leq 0$] и первый момент x фиксирован и равен a , то

$$a = \int_0^{\infty} x p(x) dx,$$

а энтропия будет максимальна при

$$p(x) = \frac{1}{a} e^{-\frac{x}{a}}$$

и этот максимум равен $\log ea$.

8. Имеется одно важное различие между энтропией непрерывных и дискретных величин. В дискретном случае энтропия изме-

рывает абсолютным образом степень случайности значения рассматриваемой случайной величины. В непрерывном случае это измерение производится *относительно заданной системы координат*. Если изменить координаты, то энтропия, вообще говоря, изменится. Действительно, при переходе к координатам $y_1 \dots y_n$ новое значение энтропии задается как

$$H(y) = \int \dots \int p(x_1, \dots, x_n) J\left(\frac{x}{y}\right) \log p(x_1, \dots, x_n) J\left(\frac{x}{y}\right) dy_1, \dots, dy_n,$$

где $J\left(\frac{x}{y}\right)$ — якобиан преобразования координат. Разлагая логарифм и заменяя переменные на $x_1, \dots, x_n$, получим

$$H(y) = H(x) - \int \dots \int p(x_1, \dots, x_n) \log J\left(\frac{x}{y}\right) dx_1, \dots, dx_n.$$

Таким образом, новое значение энтропии равно старому минус среднее значение логарифма якобиана. В непрерывном случае энтропия может рассматриваться как мера случайности *относительно принятого стандарта*, а именно выбранной системы координат, в которой каждому малому элементу объема $dx_1, \dots, dx_n$ придается одинаковый вес. Когда изменяется система координат, энтропия в новой системе также является мерой случайности, но теперь придается одинаковый вес равным элементам объема $dy_1, \dots, dy_n$ в новой системе координат.

Несмотря на зависимость от системы координат, понятие энтропии является столь же важным в непрерывном случае, как и в дискретном. Это объясняется тем, что скорость создания сообщения и пропускная способность канала определяются *разностью* двух энтропий, а эта разность *не зависит* от системы координат, так как каждая из этих двух величин изменяется на одно и то же число.

Энтропия непрерывного распределения может быть отрицательна. Масштабом измерений устанавливается произвольный нуль, соответствующий равномерному распределению по единичному объему. Распределение, более сосредоточенное чем это, будет иметь меньшую энтропию и, следовательно, будет отрицательно. Однако скорость создания сообщения и пропускная способность канала всегда будут неотрицательны.

9. Частным случаем изменения координат является линейное преобразование

$$y_j = \sum_i a_{ij} x_i.$$

В этом случае якобиан равен просто определителю $|a_{ij}|^{-1}$ и

$$H(y) = H(x) + \log |a_{ij}|.$$

В случае поворота координат (или любого сохраняющего меру преобразования) $J = 1$ и $H(y) = H(x)$.

21. Энтропия ансамбля функций

Рассмотрим эргодический ансамбль функций с ограниченной полосой частот в W герц. Пусть

$$p(x_1, \dots, x_n)$$

плотности распределения вероятностей для амплитуд $x_1, \dots, x_n$ в n последовательных выборочных точках. Определим энтропию ансамбля на степень свободы как

$$H' = - \lim_{n \rightarrow \infty} \frac{1}{n} \int \dots \int p(x_1, \dots, x_n) \log p(x_1, \dots, x_n) dx_1, \dots, dx_n.$$

Можно также определить энтропию H в секундах, проводя деление не на n , а на время T в секундах, требуемое для получения n выборочных точек. Так как $n = 2TW$, то $H' = 2WH$.

Для белого теплового шума плотность p является гауссовской и имеем

$$H' = \log \sqrt{2\pi e N},$$

$$H = W \log 2\pi e N.$$

При данной средней мощности N белый шум имеет максимальную возможную энтропию. Это следует из отмеченных выше свойств максимальной нормального распределения.

Энтропия непрерывного вероятностного процесса имеет много свойств, аналогичных свойствам энтропии дискретных процессов. В дискретном случае энтропия была связана с логарифмом *вероятности* длинных последовательностей и с *числом* относительно вероятных последовательностей большой длительности. В непрерывном случае энтропия подобным же образом связана с логарифмом *плотности вероятности* для длинного ряда выборок и с *объемом* области сравнительно высокой вероятности в функциональном пространстве.

Более точно, если предположить, что $p(x_1, \dots, x_n)$ непрерывна по всем x_i для всех n , то для достаточно большого n

$$\left| \frac{\log p}{n} - H' \right| < \epsilon$$

при любом выборе значений $(x_1, \dots, x_n)$, не принадлежащих множеству, полная вероятность которого меньше чем δ , где ϵ и δ произвольно малы. Это следует из эргодического свойства, если мы разделим пространство на большое число малых ячеек.

Связь H с объемом может быть установлена следующим образом. При тех же самых предположениях рассмотрим n -мерное пространство, соответствующее $p(x_1, \dots, x_n)$. Пусть $V_n(q)$ — наименьший объем области в этом пространстве, имеющей полную вероятность q . Тогда

$$\lim_{n \rightarrow \infty} \frac{\log V_n(q)}{n} = H',$$

если только q не равно 0 или 1.

Из сказанного видно, что при больших n существует довольно четко определенная (по крайней мере в логарифмическом смысле) область высоких вероятностей и что внутри этой области плотность распределения вероятностей относительно равномерна (опять-таки в логарифмическом смысле).

В случае белого шума распределение задается выражением

$$p(x_1, \dots, x_n) = \frac{1}{(2\pi N)^{n/2}} \exp \left\{ -\frac{1}{2N} \sum_{i=1}^n x_i^2 \right\}.$$

Так как эта функция зависит только от $\sum x_i^2$, то поверхности равной плотности распределения представляют собой сферы и все распределение обладает сферической симметрией. Областью высокой вероятности является сфера радиуса $\sqrt{nN}$. При $n \rightarrow \infty$ вероятность нахождения вне сферы радиуса $\sqrt{n(N + \epsilon)}$ стремится к нулю, как бы ни было мало ϵ , а $\frac{1}{n}$ логарифма объема сферы стремится к $\log \sqrt{2\pi e N}$.

В непрерывном случае удобно пользоваться не энтропией ансамбля H , а производной величиной, которую мы назовем *энтропийной мощностью*. Она определяется как мощность белого шума, ограниченного той же полосой частот, что и первоначальный ансамбль, и имеющего ту же самую энтропию. Другими словами, если H' — энтропия ансамбля, то его энтропийная мощность равна

$$N_1 = \frac{1}{2\pi e} \exp 2H'.$$

В геометрической трактовке это означает измерение объема высокой вероятности квадратом радиуса сферы, имеющей такой же объем. Так как белый шум имеет максимальную энтропию при данной мощности, то энтропийная мощность любого шума меньше или равна его действительной мощности.

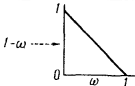
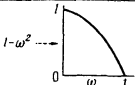
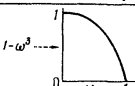
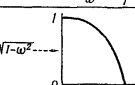
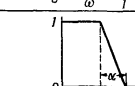
22. Потеря энтропии в линейных фильтрах

Теорема 14. Если ансамбль функций, имеющий энтропию на степень свободы H_1 в полосе частот W , пропускается через фильтр

с характеристикой $Y(f)$, то ансамбль на выходе имеет энтропию

$$H_2 = H_1 + \frac{1}{W} \int_W \log |Y(f)|^2 df.$$

Действие фильтра представляет собой линейное преобразование координат. Если рассматривать различные частотные компоненты

Усиление	Коэффициент энтропийной мощности	Усиление энтропийной мощности в децибелах	Импульсный отклик
	$\frac{1}{e^2}$	-8,68	$\frac{\sin^2 \pi t}{(\pi t)^2}$
	$\left(\frac{2}{e}\right)^4$	-5,32	$2 \left[\frac{\sin t}{t^3} - \frac{\cos t}{t^2} \right]$
	0,384	-4,15	$6 \left[\frac{\cos t - 1}{t^4} - \frac{\cos t}{2t^2} + \frac{\sin t}{t^3} \right]$
	$\left(\frac{2}{e}\right)^2$	-2,66	$\frac{\pi}{2} \frac{J_1(t)}{t}$
	$\frac{1}{e^{2\alpha}}$	-8,68 α	$\frac{1}{\alpha t^2} [\cos(1-\alpha)t - \cos t]$

как первоначальные координаты системы, то новые частотные компоненты будут просто равны исходным, умноженным на некоторые коэффициенты. Поэтому матрица преобразования координат является по существу диагональной в этих координатах. Якобиан

преобразования равен (для n синусоидальных и n косинусоидальных компонент)

$$J = \prod_{i=1}^n |Y(f_i)|^2 = \exp \sum \log |Y(f_i)|^2,$$

где f_i расположены на равных расстояниях в полосе частот W . В пределе это выражение превращается в

$$\exp \frac{1}{W} \int_W \log |Y(f)|^2 df.$$

Так как J — константа, то ее среднее значение равно ей самой, и, применяя теорему об изменении энтропии с изменением координат, получаем сформулированный выше результат. Можно выразить его в терминах энтропийной мощности. Таким образом, если энтропийная мощность первого ансамбля есть N_1 , то энтропийная мощность второго ансамбля равна

$$N_1 \exp \frac{1}{W} \int_W \log |Y(f)|^2 df.$$

Конечная энтропийная мощность равна начальной, умноженной на средне-геометрическое усиление фильтра. Если это усиление измеряется в децибелах, то выходная энтропийная мощность увеличится на средне-арифметическое усиление фильтра в децибелах в полосе частот W .

В таблице на стр. 302 потеря энтропийной мощности сосчитана (и выражена в дб) для нескольких идеализированных характеристик усиления. Приведены также импульсные отклики этих фильтров для $W = 2\pi$, причем предполагается, что фаза равна нулю.

Потеря энтропии для многих других случаев может быть найдена при помощи этих рассуждений. Например, коэффициент энтропийной мощности $\frac{1}{e^2}$, полученный для первого случая, применим также к любой характеристике усиления, получаемой из $1-\omega$ с помощью преобразования оси ω , сохраняющего меру. В частности, линейно возрастающее усиление $G(\omega) = \omega$ или пилообразная характеристика между 0 и 1 имеют такую же потерю энтропии. Обратная характеристика имеет обратный коэффициент. Поэтому $\frac{1}{\omega}$ имеет коэффициент e^2 . Возведение усиления в какую-нибудь степень приводит к возведению коэффициента в эту степень.

23. Энтропия суммы двух ансамблей

Если имеются два ансамбля функций $f_a(t)$ и $g_b(t)$, то можно образовать новый ансамбль путем «сложения». Допустим, что первый ансамбль имеет плотность распределения вероятностей $p(x_1, \dots, x_n)$,

а второй — $q(x_1, \dots, x_n)$. Тогда плотность распределения для суммы дается сверткой:

$$r(x_1, \dots, x_n) = \int \dots \int p(y_1, \dots, y_n) q(x_1 - y_1, \dots, x_n - y_n) dy_1, \dots, dy_n.$$

Физически это соответствует сложению шумов или сигналов, представляемых первоначальными ансамблями функций¹⁾.

Следующее положение доказывается в приложении 6.

Теорема 15. Пусть средние мощности двух ансамблей равны N_1 и N_2 , а их энтропийные мощности пусть равны $\bar{N}_1$ и $\bar{N}_2$. Тогда энтропийная мощность суммы, $\bar{N}_3$, ограничена неравенствами

$$\bar{N}_1 + \bar{N}_2 \leq \bar{N}_3 \leq N_1 + N_2.$$

Белый гауссовский шум обладает характерным свойством поглощать любой другой шум или ансамбль сигналов, которые могут быть сложены с ним, и при этом результирующая энтропийная мощность приблизительно равна сумме мощности белого шума и мощности сигнала (измеренной от среднего значения сигнала, которое обычно равно нулю), если только мощность сигнала мала (в некотором смысле) по сравнению с шумом.

Рассмотрим функциональное пространство, связанное с этими ансамблями и имеющее размерность n . Белый шум соответствует сферическому гауссовскому распределению в этом пространстве. Ансамбль сигналов соответствует другому распределению, не обязательно нормальному или сферическому. Пусть вторые моменты этого распределения относительно его центра тяжести равны a_{ij} . Другими словами, если $p(x_1, \dots, x_n)$ — плотность распределения вероятностей, то

$$a_{ij} = \int \dots \int p(x_1, \dots, x_n) (x_i - \alpha_i) (x_j - \alpha_j) dx_1, \dots, dx_n,$$

где α_i — координаты центра тяжести. Далее, a_{ij} является положительно определенной квадратичной формой, и можно повернуть нашу систему координат так, чтобы направить оси координат по главным направлениям этой формы. Тогда a_{ij} приведет к диагональной форме b_{ii} . Потребуем, чтобы каждое b_{ii} было мало по сравнению с N — квадратом радиуса сферического распределения.

В этом случае свертка шума и сигнала дает приблизительно гауссовское распределение, соответствующая квадратичная форма которого есть

$$N + b_{ii}.$$

¹⁾ Здесь подразумевается, конечно, что ансамбли статистически независимы. — Прим. ред.

Энтропийная мощность этого распределения равна

$$\left[\prod_{i=1}^n (N + b_{ii}) \right]^{\frac{1}{n}}$$

или приблизительно

$$\left(N^n + \sum_{i=1}^n b_{ii} N^{n-1} \right)^{\frac{1}{n}} \approx N + \frac{1}{n} \sum_{i=1}^n b_{ii}.$$

Последнее слагаемое есть мощность сигнала, первое — мощность шума.

IV. НЕПРЕРЫВНЫЙ КАНАЛ

24. Пропускная способность непрерывного канала

В непрерывном канале входные или передаваемые сигналы будут непрерывными функциями времени $f(t)$, принадлежащими к некоторому множеству, а выходные или принимаемые сигналы будут их искаженными вариантами. Рассмотрим только такой случай, когда как передаваемые, так и принимаемые сигналы ограничены некоторой полосой частот W . Тогда в интервале T они могут быть заданы с помощью $2TW$ чисел, а их статистическая структура описана конечномерными функциями распределения. Таким образом, статистика передаваемого сигнала будет определяться функцией

$$P(x_1, \dots, x_n) = P(x),$$

а статистика шума — условным распределением вероятностей

$$P_{x_1, \dots, x_n}(y_1, \dots, y_n) = P_x(y).$$

Скорость передачи информации для непрерывного канала определяется аналогично скорости передачи информации для дискретного канала, а именно

$$R = H(x) - H_y(x),$$

где $H(x)$ — энтропия входа, а $H_y(x)$ — ненадежность. Пропускная способность канала C определяется как максимум R при варьировании входа по всем возможным ансамблям. Это означает, что для конечномерного приближения надо варьировать $P(x) = P(x_1, \dots, x_n)$ и находить максимум

$$- \int P(x) \log P(x) dx + \iint P(x, y) \log \frac{P(x, y)}{P(y)} dx dy.$$

Это можно переписать в виде

$$\iint P(x, y) \log \frac{P(x, y)}{P(x)P(y)} dx dy,$$

пользуясь тем обстоятельством, что $\int \int P(x, y) \log P(x) dx dy = \int P(x) \log P(x) dx$. Таким образом, пропускная способность канала выражается следующим образом:

$$C = \lim_{T \rightarrow \infty} \max_{P(x)} \frac{1}{T} \int \int P(x, y) \log \frac{P(x, y)}{P(x)P(y)} dx dy.$$

Из этого выражения видно, что R и C не зависят от системы координат, так как и числитель и знаменатель в $\log \frac{P(x, y)}{P(x)P(y)}$ будут умножаться на один и тот же множитель при любом взаимоднозначном преобразовании x и y . Это выражение для C в виде интеграла является более общим, чем выражение $H(x) - H_y(x)$. При надлежащей интерпретации (см. приложение 7) оно всегда будет иметь смысл, в то время как выражение $H(x) - H_y(x)$ может оказаться в некоторых случаях неопределенностью вида $\infty - \infty$. Это происходит, например, если в случае n -мерной аппроксимации x полностью сосредоточено на поверхности меньшего числа измерений.

Если основание логарифмов, используемое при вычислении $H(x)$ и $H_y(x)$, равно двум, то, как и в дискретном случае, C есть максимальное число битов, которое можно передать за одну секунду по каналу с произвольно малой ненадежностью. Это можно понять физически, разделив пространство сигналов на большое число малых ячеек. Ячейки делаются настолько малыми, чтобы плотность вероятности того, что сигнал x в результате действия шума перейдет к точке y , т. е. $P_x(y)$, была фактически постоянной по всей ячейке (как по x , так и по y). Если ячейки рассматриваются как отдельные точки, то по существу имеется дискретный канал, и использованные там доказательства будут применимы и здесь. Но физически ясно, что такая квантизация объема на отдельные точки не может в любом практическом случае существенно изменить конечный результат, если только ячейки достаточно малы. Поэтому пропускная способность будет пределом пропускных способностей для дискретных подразбиений, а это и есть пропускная способность для непрерывного случая, как она определена выше.

Математически можно показать прежде всего (см. приложение 7), что если u есть сообщение, x — передаваемый сигнал, y — принятый сигнал (искаженный шумом) и v — восстановленное по принятому сигналу сообщение, то

$$H(x) - H_y(x) \geq H(u) - H_v(u)$$

независимо от того, какие операции производились над u , чтобы получить x , или над y для получения v . Таким образом, независимо от того, как закодировать двоичные знаки для создания сигнала

или как будет декодироваться принятый сигнал для восстановления сообщения, скорость дискретной передачи двоичных знаков не превысит определенную выше пропускную способность канала. С другой стороны, при весьма общих условиях можно найти систему кодирования, обеспечивающую передачу двоичных знаков со скоростью C при сколь угодно малой ненадежности или частоте ошибок. Это справедливо, например, когда берется для сигнальных функций конечномерное аппроксимирующее пространство и при этом $P(x, y)$ оказывается непрерывной как по x , так и по y , за исключением множества точек нулевой вероятности.

Важный частный случай имеет место, когда шум прибавляется к сигналу, причем шум не зависит от сигнала (в вероятностном смысле). Тогда $P_x(y)$ есть функция только разности (векторной) $n = (y - x)$,

$$P_x(y) = Q(y - x)$$

и можно приписать шуму определенную энтропию (независимо от статистики сигнала), а именно энтропию распределения $Q(n)$. Эта энтропия будет обозначаться $H(n)$.

Теорема 16. Если сигнал и шум независимы, а принимаемый сигнал является суммой передаваемого сигнала и шума, то скорость передачи равна

$$R = H(y) - H(n),$$

т. е. энтропии принимаемого сигнала за вычетом энтропии шума. Пропускная способность канала равна

$$C = \max_{P(x)} H(y) - H(n).$$

Так как $y = x + n$, то имеем

$$H(x, y) = H(x, n).$$

Разлагая левую часть и пользуясь независимостью x и n , найдем

$$H(y) + H_y(x) = H(x) + H(n).$$

Отсюда

$$\S R = H(x) - H_y(x) = H(y) - H(n).$$

Так как $H(n)$ не зависит от $P(x)$, то для максимизации R нужно максимизировать $H(y)$ — энтропию принимаемого сигнала. Если на ансамбль передаваемых сигналов наложены некоторые ограничения, то энтропия принимаемых сигналов должна быть максимизирована с учетом этих же ограничений.

25. Пропускная способность канала при ограничении средней мощности

Теорему 16 особенно просто применить для того случая, когда шум является белым тепловым шумом, а передаваемые сигналы ограничены по средней мощности величиной P . Тогда принимаемые сигналы имеют среднюю мощность $P + N$, где N — средняя мощность шума. Энтропия принимаемых сигналов будет максимальной тогда, когда они составят ансамбль белого шума, так как это даст наибольшую возможную энтропию при мощности $P + N$. Такая энтропия может быть получена путем подходящего выбора ансамбля передаваемых сигналов, а именно в том случае, когда они образуют ансамбль белого шума мощности P . Энтропия (в секунду) ансамбля принимаемых сигналов будет тогда равна

$$H(y) = W \log 2\pi e (P + N),$$

а энтропия шума

$$H(n) = W \log 2\pi e N.$$

Пропускная способность канала равна

$$C = H(y) - H(n) = W \log \frac{P+N}{N}.$$

Суммируя вышесказанное, получаем следующую теорему:

Теорема 17. *Пропускная способность канала с полосой частот W , в котором имеется белый тепловой шум мощности N при условии, что средняя мощность передаваемых сигналов ограничена величиной P , равна*

$$C = W \log \frac{P+N}{N}.$$

Это означает, что с помощью достаточно сложных систем кодирования можно передавать двоичные знаки со скоростью $W \log_2 \frac{P+N}{N}$ бит в секунду, при сколь угодно малой частоте ошибок. Невозможно передавать с большей скоростью при любой системе кодирования без того, чтобы частота ошибок не была бы положительна.

Для достижения этой предельной скорости передачи передаваемые сигналы должны приближаться по своим статистическим свойствам к белому шуму¹⁾. Одна система, для которой скорость подачи приближается к идеальной, может быть описана следующим образом:

Пусть созданы $M = 2^s$ выборок белого шума, каждая длительности T . Им приписываются двоичные числа от 0 до $(M - 1)$. В пере-

¹⁾ Эта и другие особенности, имеющие место для белого шума, рассматриваются с геометрической точки зрения в работе автора, см. стр. 433.

датчике последовательности сообщений разбиваются на группы по s двоичных знаков и для каждой группы в качестве сигнала передается соответствующая выборка шума. На приемном конце эти M выборок известны, и действительно принятый сигнал (искаженный шумом) сравнивается с каждой из них. Выборка, которая имеет наименьшее средне-квадратичное отклонение от принятого сигнала, принимается за переданный сигнал, по которому восстанавливается соответствующее двоичное число. Этот прием эквивалентен выбору наиболее вероятного (*апостериори*) сигнала. Число используемых выборок шума M будет зависеть от допустимой частоты ошибок ε , но для почти всех наборов выборок имеем

$$\lim_{\varepsilon \rightarrow 0} \lim_{T \rightarrow \infty} \frac{\log M(\varepsilon, T)}{T} = W \log \frac{P+N}{N}.$$

Таким образом, независимо от того, насколько малым выбрано ε , можно, выбирая T достаточно большим, приблизиться сколь угодно близко к передаче $TW \log \frac{P+N}{N}$ двоичных единиц за время T .

Формулы, подобные $C = W \log (P+N)/N$, для случая белого шума были получены независимо и некоторыми другими авторами, хотя в несколько другой интерпретации. Здесь можно упомянуть в этой связи работы Н. Винера, В. Г. Таллера и Х. Сулливана¹).

В случае произвольного искажающего шума (не обязательно белого теплового шума) задача максимизации, связанная с определением пропускной способности C , по-видимому, не может быть решена явно. Однако могут быть установлены верхняя и нижняя границы для C в терминах средней мощности шума N и энтропийной мощности шума N_1 . В большинстве практических случаев эти границы достаточно близки друг к другу и поэтому дают удовлетворительное решение проблемы.

Теорема 18. *Пропускная способность канала с полосой частот W , в котором имеется произвольный шум, ограничена неравенствами*

$$W \log \frac{P+N_1}{N_1} \leq C \leq W \log \frac{P+N}{N_1},$$

где P — средняя мощность передаваемых сигналов, N — средняя мощность шума, N_1 — энтропийная мощность шума.

Здесь опять средняя мощность искаженных сигналов будет $P+N$. Энтропия была бы максимальной для этой мощности, если

¹) Wiener N., Cybernetics (русский перевод см. сноску на стр. 294. — Прим. ред.); Tuller W., Sullivan H., Theoretical limitations on the rate of transmission of information, *PIRE*, 37, № 5, May (1949), 468.

бы принимаемый сигнал был бы белым шумом, и она равнялась бы $W \log 2\pi e(P + N)$. Может быть, этого и нельзя достичь, т. е., может быть, и не существует такого ансамбля передаваемых сигналов, который, будучи добавлен к искажающему шуму, дал бы белый тепловой шум в приемнике, но по крайней мере это устанавливает верхнюю границу для $H(y)$. Поэтому имеем

$$C = \max H(y) - H(n) \leq W \log 2\pi e(P + N) - W \log 2\pi eN_1.$$

Это и есть верхняя граница, даваемая теоремой. Нижняя граница может быть получена при рассмотрении скорости передачи для случая, когда передаваемый сигнал является белым шумом мощности P . При этом энтропийная мощность принимаемого сигнала должна быть не меньше энтропийной мощности белого шума, мощность которого равна $P + N_1$, так как в теореме 15 доказано, что энтропийная мощность суммы двух ансамблей больше или равна сумме отдельных энтропийных мощностей. Поэтому

$$\max H(y) \geq W \log 2\pi e(P + N_1)$$

и

$$C \geq W \log 2\pi e(P + N_1) - W \log 2\pi eN_1 = W \log \frac{P + N_1}{N_1}.$$

С увеличением P верхняя и нижняя границы, даваемые теоремой 18, сближаются, так что имеем в качестве асимптотической скорости

$$W \log \frac{P + N}{N_1}.$$

Если сам шум является белым, то $N = N_1$ и полученный результат сводится к формуле, доказанной ранее:

$$C = W \log \left(1 + \frac{P}{N} \right).$$

Если шум является гауссовским, но спектр его не обязательно равномерный, то N_1 является средне-геометрической мощностью шума по различным частотам в полосе W . Поэтому

$$N_1 = \exp \frac{1}{W} \int_W \log N(f) df,$$

где $N(f)$ — мощность шума на частоте f .

Теорема 19. Если при данной мощности передаваемых сигналов P пропускную способность канала обозначить через C , где

$$C = W \log \frac{P + N - \eta}{N_1},$$

то окажется, что η монотонно убывает с ростом P , приближаясь в пределе к нулю.

Предположим, что при данной мощности P_1 пропускная способность канала равна

$$W \log \frac{P_1 + N - \eta_1}{N_1}.$$

Это означает, что сигнал с наилучшим распределением, скажем $p(x)$, будучи добавлен к шуму с распределением $q(x)$, даст для принимаемого сигнала распределение $r(y)$, энтропийная мощность которого равна $(P_1 + N - \eta_1)$. Пусть мощность увеличена до $P_1 + \Delta P$ путем добавления к сигналу белого шума мощности ΔP . Энтропия принимаемого сигнала теперь равна по меньшей мере

$$H(y) = W \log 2\pi e (P_1 + N - \eta_1 + \Delta P),$$

что следует из применения теоремы о минимуме энтропийной мощности суммы двух ансамблей (теорема 15). Следовательно, так как можно достичь указанной величины H , то энтропия максимизирующего распределения должна быть не меньше этой величины, а значит, η должно монотонно убывать. Чтобы показать, что при $P \rightarrow \infty$ величина $\eta \rightarrow 0$, рассмотрим сигнал, являющийся белым шумом мощности P . Каков бы ни был искажающий шум, если P достаточно велико, то принимаемый сигнал будет приблизительно белым шумом в том смысле, что его энтропийная мощность будет стремиться к $P + N$.

26. Пропускная способность канала при ограничении пиковой мощности

В некоторых приложениях ограниченной является не средняя мощность передатчика, а его мгновенная пиковая мощность. Задача вычисления пропускной способности канала сводится тогда к максимизации (с помощью варьирования ансамбля передаваемых символов) выражения

$$H(y) - H(n)$$

при условии, что все функции ансамбля $f(t)$ при всех t меньше или равны, скажем $\sqrt{S}$. С ограничениями такого типа работать математически не так легко, как с ограничениями на среднюю мощность. Наибольшее, что удалось получить для этого случая, — это нижнюю границу для всех $\frac{S}{N}$, «асимптотическую» верхнюю границу (справедливую при больших $\frac{S}{N}$) и асимптотическое значение C для малых $\frac{S}{N}$.

Теорема 20. Пропускная способность канала C с полосой частот W , в котором передаваемые сигналы искажаются белым теп-

ловым шумом мощности N , ограничена неравенством

$$C \geq W \log \frac{2}{\pi e^3} \frac{S}{N},$$

где S — допустимая пиковая мощность в передатчике. Для достаточно больших $\frac{S}{N}$

$$C \leq W \log \frac{\frac{2}{\pi e} S + N}{N} (1 + \epsilon),$$

где ϵ стремится к нулю при $\frac{S}{N} \rightarrow \infty$. При $\frac{S}{N} \rightarrow 0$ (и при условии, что полоса частот W начинается от нуля)

$$C \rightarrow W \log \left(1 + \frac{S}{N} \right).$$

Желательно максимизировать энтропию принимаемых сигналов. Если $\frac{S}{N}$ велико, то это будет приблизительно осуществляться тогда, когда энтропия передаваемого ансамбля достигает максимума.

Асимптотическая верхняя граница получается с помощью ослабления условий, наложенных на ансамбль. Предположим, что мощность ограничена величиной S не в каждый момент времени, а лишь в выборочных точках. Максимум энтропии передаваемого ансамбля при таких ослабленных условиях будет больше или равен максимуму при исходных условиях. Измененная таким образом задача может быть легко решена. Энтропия будет максимальной, если различные выборки независимы и имеют равномерное распределение на интервале $-\sqrt{S}$, $+\sqrt{S}$. Энтропия при этом равна $W \log 4S$.

Принимаемые сигналы будут тогда иметь энтропию, меньшую чем $W \log (4S + 2\pi eN) \cdot (1 + \epsilon)$,

причем $\epsilon \rightarrow 0$ при $\frac{S}{N} \rightarrow \infty$ и пропускная способность канала получается путем вычитания из этого выражения энтропии белого шума $W \log 2\pi eN$:

$$W \log (4S + 2\pi eN) (1 + \epsilon) - W \log (2\pi eN) = W \log \frac{\frac{2}{\pi e} S + N}{N} (1 + \epsilon).$$

Это и есть искомая верхняя граница для пропускной способности канала.

Чтобы получить нижнюю границу, рассмотрим тот же самый ансамбль функций. Пусть эти функции пропускаются через идеаль-

ный линейный фильтр с треугольной характеристикой. Усиление должно быть единичным на нулевой частоте и линейно спадать до нуля на частоте W . Покажем сначала, что функции на выходе фильтра имеют пиковое ограничение S во все моменты времени (а не только в выборочных точках). Заметим прежде всего, что импульс $\frac{\sin 2\pi Wt}{2\pi Wt}$, проходя через фильтр, дает на выходе

$$\frac{1}{2} \frac{\sin^2 \pi Wt}{(\pi Wt)^2}.$$

Эта функция всегда неотрицательна. Входная функция в общем случае может рассматриваться как сумма ряда сдвинутых функций

$$a \frac{\sin 2\pi Wt}{2\pi Wt},$$

где амплитуда в момент выбора a не превосходит $\sqrt{S}$. Поэтому выход представляет собой сумму сдвинутых неотрицательных функций указанного типа с теми же коэффициентами¹⁾. Так как эти функции неотрицательны, то их наибольшее положительное значение для любого момента времени t получается тогда, когда все коэффициенты a имеют максимальную положительную величину, т. е. $\sqrt{S}$. В этом случае входная функция представляет собой константу амплитуды $\sqrt{S}$, а так как фильтр имеет единичное усиление для постоянной составляющей, то выход будет тот же самый. Поэтому выходной ансамбль имеет пиковую мощность S .

Энтропия выходного ансамбля может быть вычислена через энтропии входного ансамбля при помощи доказанной ранее теоремы. Выходная энтропия равна входной энтропии плюс средне-геометрическое усиление фильтра

$$\int_0^W \log G^2 df = \int_0^W \log \left(\frac{W-f}{W} \right)^2 df = -2W.$$

Следовательно, выходная энтропия равна

$$W \log 4S - 2W = W \log \frac{4S}{e^2},$$

и пропускная способность канала больше чем

$$W \log \frac{2}{\pi e^3} \frac{S}{N}.$$

Теперь требуется показать, что для малых S/N (отношение пиковой мощности сигнала к средней мощности белого шума) про-

¹⁾ На самом деле эти коэффициенты умножаются на $1/2$, что несущественно для дальнейшего рассуждения. — *Прим ред.*

пускная способность канала равна

$$C = W \log \left(1 + \frac{S}{N} \right).$$

Более точно $C/W \log \left(1 + \frac{S}{N} \right) \rightarrow 1$ при $\frac{S}{N} \rightarrow 0$. Так как средняя мощность сигнала P меньше или равна пиковой мощности S , то отсюда следует, что для всех S/N

$$C \leq W \log \left(1 + \frac{P}{N} \right) \leq W \log \left(1 + \frac{S}{N} \right).$$

Поэтому, если удастся найти ансамбль функций, соответствующих скорости передачи, близкой к $W \log(1 + S/N)$, и ограниченных полосой частот W и пиковой мощностью S , то последняя часть теоремы будет доказана. Рассмотрим ансамбль функций следующего типа: пусть t функций в последовательные моменты выбора принимают одинаковое значение (либо $+\sqrt{S}$, либо $-\sqrt{S}$), в следующие t моментов выбора опять имеют одинаковое значение и т. д. Знак для последовательности выбирается случайным образом: с вероятностью $1/2$ берется $+\sqrt{S}$ и с вероятностью $1/2$ берется $-\sqrt{S}$. Если этот ансамбль пропустить через фильтр с треугольной характеристикой усиления и равным единице усилением для постоянной составляющей, то выходной сигнал имеет пиковую мощность, не превосходящую $+S$. Кроме того, средняя мощность близка к S и приближается к этому значению при увеличении t . Энтропия суммы этого ансамбля и теплового шума может быть найдена с помощью теоремы о сумме шума и малого сигнала. Эта теорема применима, если выражение

$$\sqrt{t} \frac{S}{N}$$

достаточно мало. Это можно обеспечить, взяв отношение $\frac{S}{N}$ достаточно малым (после того, как выбрано t). Энтропийная мощность со сколь угодно близким приближением будет равна $S + N$, и, следовательно, скорость передачи может быть сделана сколь угодно близкой к

$$W \log \left(\frac{S+N}{N} \right).$$

V. СКОРОСТЬ СОЗДАНИЯ СООБЩЕНИЙ ДЛЯ НЕПРЕРЫВНОГО ИСТОЧНИКА

27. Функции оценки точности воспроизведения

В случае дискретного источника сообщений было можно установить определенную скорость создания сообщений, а именно энтропию соответствующего вероятностного процесса. Для непрерыв-

ного источника положение оказывается значительно сложнее. Прежде всего непрерывно изменяющаяся величина может принимать бесконечное число значений, и поэтому для точного задания этой величины требуется бесконечное число двоичных единиц. Это означает, что при передаче от непрерывного источника для точного воспроизведения в точке приема, вообще говоря, необходим канал с бесконечной пропускной способностью (в битах в сек.). Поскольку в каналах имеется обычно некоторый уровень шумов и, следовательно, пропускная способность ограничена, точная передача невозможна.

Это рассуждение, однако, обходит основной вопрос. Практически при непрерывном источнике нас интересует не точная передача, а передача с определенным допуском. Вопрос заключается в том, можно ли приписать непрерывному источнику определенную скорость создания сообщений в том случае, когда требуется только определенная точность воспроизведения, измеренная подходящим способом. Разумеется, при возрастании требований к точности воспроизведения скорость будет возрастать. Как будет показано ниже, можно в очень общих случаях определить такую скорость создания сообщений. Она будет обладать тем свойством, что при надлежащем кодировании можно передать по каналу информацию, удовлетворив при этом требования к точности воспроизведения, если только пропускная способность канала равна рассматриваемой скорости. Меньшая пропускная способность оказывается для этого недостаточной.

Прежде всего необходимо дать общую математическую формулировку понятия точности передачи. Рассмотрим множество сообщений большой длительности, скажем T секунд. Источник описывается заданием в соответствующем пространстве плотности распределения вероятностей $P(x)$ того, что будет выбрано данное сообщение. Данная система связи описывается (с внешней точки зрения) заданием условной вероятности $P_x(y)$ того, что если источник создал сообщение x , то воспроизводимое сообщение в точке приема будет y .

Система в целом (включая источник и передающую систему) описывается совместным распределением вероятностей $P(x, y)$ того, что имеются передаваемое сообщение x и воспроизведенное сообщение y . Если известна эта функция, то тем самым полностью известны свойства системы с точки зрения точности воспроизведения. Любая оценка точности воспроизведения должна математически соответствовать некоторой операции над функцией $P(x, y)$. Эта операция должна по крайней мере приводить к простому упорядочению систем, т. е. для двух систем, задаваемых функциями $P_1(x, y)$ и $P_2(x, y)$, эта операция должна давать в соответствии с нашим критерием точности один из трех ответов: 1) первая система дает более высокую точность, 2) вторая система дает более высокую

точность или же 3) они дают одинаковую точность. Это означает, что критерий точности может быть представлен принимающей числовые значения функцией оценки:

$$v[P(x, y)],$$

аргументом которой может быть любое возможное распределение вероятностей $P(x, y)$. Функция $v[P(x, y)]$ упорядочивает системы связи соответственно их точности воспроизведения, и для удобства примем, что меньшим значениям v соответствует «более высокая точность».

Теперь покажем, что при самых общих и разумных предположениях функция $v[P(x, y)]$ может быть записана в форме, кажущейся много более частной, а именно как среднее некоторой функции $q(x, y)$, взятое по множеству возможных значений x и y :

$$v[P(x, y)] = \int \int P(x, y) q(x, y) dx dy.$$

Для того чтобы это показать, достаточно предположить, 1) что источник и система являются эргодическими, так что очень длинная выборка будет с вероятностью, близкой к единице, типичной для ансамбля, и 2) что оценка является «разумной» в том смысле, что возможно на основе наблюдения типичной входной выборки x_1 и типичной выходной выборки y_1 образовать выборочную оценку, и если длительность этих выборок возрастает, то выборочная оценка будет с вероятностью единицы стремиться к точной оценке, основанной на полном знании $P(x, y)$. Пусть выборочная оценка обозначена через $q(x, y)$. Тогда функция $q(x, y)$ стремится (при $T \rightarrow \infty$) к постоянной величине для почти всех (x, y) , которые находятся в высоковероятной области для данной системы:

$$q(x, y) \rightarrow v[P(x, y)];$$

можно также записать

$$q(x, y) \rightarrow \int \int P(x, y) q(x, y) dx dy,$$

так как

$$\int \int P(x, y) dx dy = 1,$$

что и требовалось доказать.

Функция $q(x, y)$ обладает общими свойствами «расстояния» между x и y ¹⁾. Она измеряет, насколько нежелательно (в соответствии с нашим критерием точности воспроизведения) принять y ,

¹⁾ Однако она не является «метрикой» в строгом смысле, так как, вообще говоря, не удовлетворяет условию $q(x, y) = q(y, x)$ или условию $q(x, y) + q(y, z) \geq q(x, z)$.

когда передано x . Полученный выше общий результат может быть переформулирован следующим образом: любая разумная оценка может быть представлена как среднее значение функции расстояния, усредненной по множеству исходных и воспроизводимых сообщений x и y , в соответствии с вероятностью $P(x, y)$, при условии, что длительность сообщений T берется достаточно большой.

Простыми примерами функций оценки являются следующие:

1. Средне-квадратичный критерий

$$v = \overline{[x(t) - y(t)]^2}.$$

В этом очень часто применяемом критерии точности функция расстояния $q(x, y)$ равна (с точностью до постоянного множителя) квадрату обычного евклидова расстояния между точками x и y в соответствующем пространстве

$$q(x, y) = \frac{1}{T} \int_0^T [x(t) - y(t)]^2 dt.$$

2. Частотно-взвешенный средне-квадратичный критерий. Некоторое обобщение предыдущего критерия заключается в том, что перед использованием средне-квадратичной меры точности можно приписать различные веса разным частотным компонентам. Это эквивалентно пропусканию разности $x(t) - y(t)$ через формирующий фильтр с последующим определением средней мощности на выходе.

Итак, пусть

$$e(t) = x(t) - y(t)$$

и

$$f(t) = \int_{-\infty}^{\infty} e(\tau) k(t - \tau) d\tau.$$

Тогда

$$q(x, y) = \frac{1}{T} \int_0^T f^2(t) dt.$$

3. Критерий абсолютной ошибки

$$q(x, y) = \frac{1}{T} \int_0^T |x(t) - y(t)| dt.$$

4. Строение уха и мозга определяет неявно несколько оценок, применимых для случаев передачи речи или музыки. Так, например, существует критерий «понятности», в котором $q(x, y)$ равна относительной частоте неправильно интерпретированных слов,

когда сообщение $x(t)$ принимается как $y(t)$. Хотя в этих случаях нельзя дать явного выражения для функции $q(x, y)$, она все же может быть определена в принципе из достаточно обширных экспериментов. Некоторые ее свойства следуют из хорошо известных экспериментальных результатов по исследованию слуха, например из того, что ухо сравнительно нечувствительно к фазе, а его чувствительность к амплитуде и частоте приблизительно логарифмическая.

5. Дискретный случай может рассматриваться как тот частный случай, когда молчаливо подразумевается оценка, основанная на частоте ошибок. Функция $q(x, y)$ определяется в этом случае как число символов в последовательности y , отличающихся от соответствующих символов в последовательности x , деленное на полное число символов в последовательности x .

28. Скорость создания сообщений источником при данной точности воспроизведения

Теперь можно определить скорость создания сообщений непрерывным источником. Нам задано распределение $P(x)$ для источника и оценка v , определяемая функцией расстояния $q(x, y)$, которая будет предполагаться непрерывной как по x , так и по y . Качество данной системы $P(x, y)$ измеряется величиной

$$v = \iint q(x, y) P(x, y) dx dy.$$

Кроме того, скорость потока двоичных символов, соответствующая $P(x, y)$, равна

$$R = \iint P(x, y) \log \frac{P(x, y)}{P(x)P(y)} dx dy.$$

Определим скорость создания сообщений R_1 при данной точности воспроизведения v_1 как минимум R , когда при фиксированном $v = v_1$ варьируется $P_{x(y)}$. Таким образом,

$$R_1 = \min_{P_{x(y)}} \iint P(x, y) \log \frac{P(x, y)}{P(x)P(y)} dx dy$$

при условии

$$v_1 = \iint P(x, y) q(x, y) dx dy.$$

Это означает, что в действительности рассматриваются все системы связи, которые могли бы быть использованы и которые обеспечивают передачу с требуемой точностью. Скорость передачи в битах в секунду вычисляется для каждой системы, и выбирается наименьшая скорость. Она и будет скоростью создания сообщений, которую мы приписываем источнику при данной точности воспроизведе-

дения. Обоснование этого определения заключается в следующей теореме.

Теорема 21¹⁾. Если источник при данной оценке v_1 имеет скорость создания сообщений R_1 , то можно закодировать сообщения на выходе источника и передавать их по каналу с пропускной способностью C при точности воспроизведения, как угодно близкой к v_1 , если только $R_1 \leq C$. Это невозможно, если $R_1 > C$.

Последнее утверждение теоремы немедленно следует из определения R_1 и предыдущих рассуждений. Если бы оно не было верным, то можно было бы передавать больше чем C бит в секунду по каналу с пропускной способностью C . Первая часть теоремы доказывается методом, аналогичным использованному при доказательстве теоремы 11. Во-первых, можно разделить пространство (x, y) на большое число малых ячеек и рассматривать этот случай как дискретный. Это изменит функцию оценки не больше чем на сколь угодно малую величину (если ячейки очень малы) из-за предполагаемой непрерывности $q(x, y)$. Предположим, что $P_1(x, y)$ есть конкретная система, которая минимизирует скорость, придавая ей величину R_1 . Выберем из высоковероятных сообщений случайным образом множество, содержащее

$$2^{(R_1 + \varepsilon)T}$$

членов, где $\varepsilon \rightarrow 0$ при $T \rightarrow \infty$. При большом T каждая выбранная точка будет соединяться высоковероятными линиями (как на рис. 10) с некоторым множеством x . Вычисления, подобные использованным при доказательстве теоремы 11, показывают, что при большом T почти все x охватываются веерами линий от выбранных точек y для почти всякого выбора множества y . Система связи, которая должна быть использована, действует следующим образом. Выбранным точкам приписываются двоичные числа. Когда появляется сообщение x , оно будет (с вероятностью, стремящейся к 1 при $T \rightarrow \infty$) расположено по крайней мере на одном из вееров. Соответствующее двоичное число (или если их несколько, то одно из них, выбранное произвольно) передается по каналу, закодированное надлежащим образом для обеспечения малой вероятности ошибки. Так как $R_1 \leq C$; то это возможно. В точке приема восстанавливается соответствующее y , и затем оно используется как воспроизведенное сообщение.

Оценка v'_1 для этой системы может быть сделана сколь угодно близкой к v_1 , если взять T достаточно большим. Это происходит из-за того, что для каждой длинной выборки сообщения $x(t)$ и воспроизведенного сообщения $y(t)$ оценка стремится к v_1 (с вероятностью 1).

¹⁾ Математически полные формулировки и доказательство этой теоремы можно найти в работе Р. Л. Добрушина, см. сноску на стр. 281. — *Прим. ред.*

Интересно отметить, что в этой системе шумы в воспроизводимом сообщении в действительности создаются за счет квантования в передатчике, а не за счет шума в канале. Они более или менее аналогичны шумам квантования при кодово-импульсной модуляции.

29. Вычисление скорости создания сообщений

Определение скорости создания сообщений во многих отношениях аналогично определению пропускной способности канала. В первом случае

$$R = \min_{P(x|y)} \iint P(x, y) \log \frac{P(x, y)}{P(x)P(y)} dx dy$$

при фиксированных

$$P(x) \quad \text{и} \quad v_1 = \iint P(x, y) q(x, y) dx dy.$$

Во втором случае

$$C = \max_{P(x)} \iint P(x, y) \log \frac{P(x, y)}{P(x)P(y)} dx dy$$

при фиксированном $P_x(y)$ и при наложении возможно одного или более других ограничений (например, ограничения средней мощности) вида

$$K = \iint P(x, y) \lambda(x, y) dx dy$$

Для общей задачи о нахождении максимума, возникающей при вычислении скорости создания сообщений источником, можно найти решение в частном случае. Используя метод Лагранжа, рассмотрим

$$\iint \left[P(x, y) \log \frac{P(x, y)}{P(x)P(y)} + \mu P(x, y) q(x, y) + \nu(x) P(x, y) \right] dx dy.$$

Вариационное уравнение [когда берется первая вариация по $P(x, y)$] приводит к равенству

$$P_y(x) = B(x) e^{-\lambda q(x, y)},$$

где λ определяется из условия для достижения требуемой точности воспроизведения, а $B(x)$ подбирается так, чтобы удовлетворялось равенство

$$\int B(x) e^{-\lambda q(x, y)} dx = 1.$$

Это показывает, что при наилучшем кодировании условная вероятность $P_y(x)$ экспоненциально зависит от $q(x, y)$ — функции расстояния между рассматриваемыми x и y .

В частном случае, когда функция расстояния $\varrho(x, y)$ зависит только от разности (векторной) между x и y ,

$$\varrho(x, y) = \varrho(x - y)$$

имеем

$$\int B(x) e^{-\lambda \varrho(x-y)} dx = 1.$$

Следовательно, $B(x)$ — константа, скажем α , и

$$P_y(x) = \alpha e^{-\lambda \varrho(x-y)}.$$

К сожалению, в конкретных случаях эти формальные решения трудно численно оценить и поэтому их ценность представляется небольшой. Фактическое вычисление скорости создания сообщений удалось провести лишь для немногих простых случаев.

Если функция расстояния $\varrho(x, y)$ является средне-квадратичным отклонением x от y и ансамбль сообщений представляет собой белый шум, то скорость создания сообщений может быть определена. В этом случае мы имеем

$$R = \min [H(x) - H_y(x)] = H(x) - \max H_y(x),$$

и при этом $N = \overline{(x-y)^2}$. Но максимум $H_y(x)$ достигается в том случае, когда $y-x$ есть белый шум, и этот максимум равен $W_1 \log 2\pi eN$, где W_1 — полоса частот ансамбля сообщений. Поэтому

$$R = W_1 \log 2\pi eQ - W_1 \log 2\pi eN = W_1 \log \frac{Q}{N},$$

где Q — средняя мощность сообщений. Таким образом, доказана следующая теорема:

Теорема 22. Скорость создания сообщений для источника белого шума мощности Q с полосой частот W_1 при средне-квадратичном критерии точности воспроизведения равна

$$R = W_1 \log \frac{Q}{N},$$

где N — допустимый средний квадрат отклонения воспроизводимого сообщения от исходного.

В более общем случае для любого источника сообщений можно получить неравенства, ограничивающие скорость создания сообщений при критерии средне-квадратичной ошибки.

Теорема 23. Скорость создания сообщений для любого источника с полосой частот W_1 ограничена неравенствами:

$$W_1 \log \frac{Q_1}{N} \leq R \leq W_1 \log \frac{Q}{N},$$

где Q — средняя мощность источника, Q_1 — энтропийная мощность источника и N — допустимая средне-квадратичная ошибка.

Нижняя граница следует из того, что $\max H_y(x)$ при данном $\overline{(x-y)^2} = N$ достигается для белого шума. Верхняя граница получается, если разместить точки, использованные при доказательстве теоремы 21, не лучшим образом, а случайно в сфере радиуса $\sqrt{Q-N}$.

Благодарность

Автор благодарен своим коллегам по лаборатории, особенно докторам Боде, Пирсу, Мак-Миллану, Оливеру за многие полезные предложения и критические замечания, высказанные во время работы над настоящей статьей. Следует выразить признательность также профессору Винеру, чье изящное решение задач фильтрации и предсказания стационарных ансамблей оказало значительное влияние на представления автора в этой области.

Приложение I

РОСТ ЧИСЛА БЛОКОВ СИМВОЛОВ ПРИ УСЛОВИЯХ, ОПИСЫВАЕМЫХ КОНЕЧНЫМ ЧИСЛОМ СОСТОЯНИЙ¹⁾

Пусть $N_i(L)$ — число блоков символов длины L , оканчивающихся в состоянии i . Тогда имеем

$$N_j(L) = \sum_{i,s} N_i(L - b_{ij}^{(s)}),$$

где $b_{ij}^{(1)}, b_{ij}^{(2)}, \dots, b_{ij}^{(s)}$ — длительности символов, которые могут быть выбраны в состоянии i и привести к состоянию j . Написанные равенства представляют собой линейные разностные уравнения, решения которых при $L \rightarrow \infty$ должны иметь вид

$$N_j = A_j W^L.$$

Подставим в разностное уравнение это выражение

$$A_j W^L = \sum_{i,s} A_i W^{L - b_{ij}^{(s)}}$$

¹⁾ Математически строгое доказательство теоремы 1 дано в работе Л. Юбич Ю. И., Замечание о пропускной способности дискретного канала связи без шумов, *Успехи матем. наук*, 17 (1962), № 1, 191. — *Прим. ред.*

или

$$A_j = \sum_{is} A_i W^{-b_{ij}^{(s)}},$$

$$\sum_i \left(\sum_s W^{-b_{ij}^{(s)}} - \delta_{ij} \right) A_i = 0.$$

Для того чтобы не все A_i были равны 0, необходимо, чтобы определитель

$$D(W) = |a_{ij}| = \left| \sum_s W^{-b_{ij}^{(s)}} - \delta_{ij} \right|$$

был равен нулю. Это определяет значение W , которое выбирается, конечно, как наибольший действительный корень уравнения $D = 0$.

Величина C при этом равна

$$C = \lim_{L \rightarrow \infty} \frac{\log \sum A_j W^L}{L} = \log W.$$

Заметим также, что те же самые свойства роста получаются, если потребовать, чтобы все блоки начинались в одном и том же (произвольно выбранном) состоянии.

Приложение 2

Вывод равенства $H = -\sum p_i \log p_i$

Пусть $H\left(\frac{1}{n}, \frac{1}{n}, \dots, \frac{1}{n}\right) = A(n)$. По условию (3) можно разбить выбор из s^m равновероятных возможностей на m выборов из s равновероятных возможностей в каждом выборе. При этом должно выполняться равенство

$$A(s^m) = mA(s).$$

Точно так же

$$A(t^n) = nA(t).$$

Можно взять n произвольно большим, а m найти из условия

$$s^m \leq t^n < s^{(m+1)}.$$

Таким образом, логарифмируя и деля на $n \log s$, получим

$$\frac{m}{n} \leq \frac{\log t}{\log s} \leq \frac{m}{n} + \frac{1}{n}$$

или

$$\left| \frac{m}{n} - \frac{\log t}{\log s} \right| < \varepsilon,$$

где ε — произвольно мало. Теперь из свойства монотонности $A(n)$

$$\begin{aligned} A(s^m) &\leq A(t^m) \leq A(s^{m+1}), \\ mA(s) &\leq nA(t) \leq (m+1)A(s). \end{aligned}$$

Следовательно, деля на $nA(s)$, получаем

$$\frac{m}{n} \leq \frac{A(t)}{A(s)} \leq \frac{m}{n} + \frac{1}{n}$$

или

$$\begin{aligned} \left| \frac{m}{n} - \frac{A(t)}{A(s)} \right| &< \varepsilon, \\ \left| \frac{A(t)}{A(s)} - \frac{\log t}{\log s} \right| &\leq 2\varepsilon, \\ A(t) &= -K \log(t), \end{aligned}$$

где K должно быть положительным и удовлетворять условию 2).

Теперь допустим, что имеется выбор из n возможностей с соизмеримыми вероятностями $p_i = n_i / \sum n_i$, где n_i — целые числа. Можно разбить выбор из $\sum n_i$ возможностей на выбор из возможностей с вероятностями $p_1, \dots, p_n$ с последующим равновероятным выбором из n_i возможностей, если в первом случае выбрано i . Пользуясь опять условием 3), приравняем полные неопределенности выборов из $\sum n_i$ возможностей, сосчитанные двумя способами:

$$K \log \sum n_i = H(p_1, \dots, p_n) + K \sum p_i \log n_i.$$

Следовательно,

$$\begin{aligned} H &= K [\sum p_i \log \sum n_i - \sum p_i \log n_i] = \\ &= -K \sum p_i \log \frac{n_i}{\sum n_i} = -K \sum p_i \log p_i. \end{aligned}$$

Если p_i несоизмеримы, то их можно аппроксимировать рациональными дробями и то же самое выражение для H будет иметь место из-за предположения непрерывности. Таким образом, это выражение справедливо в общем случае. Выбор коэффициента производится из соображений удобства: он определяет единицу измерений.

Приложение 3

ТЕОРЕМА ОБ ЭРГОДИЧЕСКИХ ИСТОЧНИКАХ

Предположим, что источник является эргодическим, так что применим усиленный закон больших чисел. Таким образом, число непосредственных переходов из состояния i в состояние j в последовательности большой длины N приблизительно пропорционально

вероятности нахождения в состоянии i , скажем P_i , умноженной на p_{ij} , т. е. $P_i p_{ij} N$. Если N достаточно велико, то вероятность ошибки, не превосходящей $\delta\%$, в этом случае меньше чем ε , так что все числа, кроме множества малой вероятности, заключены в пределах

$$(P_i p_{ij} \pm \delta) N.$$

Следовательно, все последовательности, за исключением сколь угодно малой доли, имеют вероятность p :

$$p = \Pi p_{ij}^{(P_i p_{ij} \pm \delta) N}$$

и $\log p/N$ ограничен соотношениями

$$\frac{\log p}{N} = \sum (P_i p_{ij} \pm \delta) \log p_{ij}$$

или

$$\left| \frac{\log p}{N} - \sum P_i p_{ij} \log p_{ij} \right| < \eta.$$

Это доказывает теорему 3.

Теорема 4 немедленно следует отсюда, если вычислить верхнюю и нижнюю границы для $n(q)$, учитывая область возможных значений p , указываемую теоремой 3.

В смешанном (не эргодическом) случае, когда

$$L = \sum p_i L_i,$$

а энтропии компонент суть $H_1 \geq H_2 \geq \dots \geq H_n$, справедлива следующая теорема:

Теорема. Пусть $\lim_{N \rightarrow \infty} \frac{\log n(q)}{N} = \Phi(q)$; тогда $\Phi(q)$ — убывающая ступенчатая функция и в интервале $\sum_1^{s-1} \alpha_i < q < \sum_1^s \alpha_i$ справедливо равенство $\Phi(q) = H_s$.

Для доказательства теорем 5 и 6 прежде всего заметим, что F_N монотонно убывает, так как увеличение N увеличивает значение условной энтропии. Простая подстановка значения $p_{B_i}(S_j)$ в формулу для F_N показывает, что

$$F_N = N G_N - (N-1) G_{N-1}.$$

Суммируя по всем N , получим

$$G_N = \frac{1}{N} \sum_{i=1}^N F_i.$$

Следовательно, $G_N \geq F_N$ и G_N монотонно убывают. Они должны также сходиться к тому же самому пределу. Применяя теорему 3, видим, что

$$\lim_{N \rightarrow \infty} G_N = H.$$

Приложение 4

МАКСИМИЗАЦИЯ СКОРОСТИ ДЛЯ СИСТЕМЫ ОГРАНИЧЕНИЯ

Допустим, что на данные последовательности символов наложен ряд ограничений, задающих систему с конечным числом состояний, которая может быть поэтому изображена графически, как на рис. 2. Пусть $l_{ij}^{(s)}$ будут длительности различных символов, которые могут появиться при переходе из состояния i в состояние j . Каково распределение вероятностей для различных состояний P_i и вероятностей $p_{ij}^{(s)}$ выбора символа s при переходе из состояния i в состояние j , для которого максимизируется скорость создания информации при данных ограничениях? Эти ограничения определяют дискретный канал, а максимальная скорость должна быть меньше или равна пропускной способности C этого канала. Действительно, если все группы большой длительности были бы равновероятны, то в результате получилась бы именно эта скорость, и если это возможно, то такая скорость была бы наилучшей. Покажем, что эта скорость может быть достигнута при надлежащем выборе P_i и P_{ij}^s . Рассматриваемая скорость равна

$$-\frac{\sum_{i, j, s} P_i p_{ij}^{(s)} \log p_{ij}^{(s)}}{\sum_{i, j, s} P_i p_{ij}^{(s)} l_{ij}^{(s)}} = \frac{N}{M}.$$

Пусть

$$p_{ij}^{(s)} = \frac{B_j}{B_i} W^{-l_{ij}^{(s)}},$$

где B_i удовлетворяют системе уравнений

$$B_i = \sum_{j, s} B_j W^{-l_{ij}^{(s)}}.$$

Эта однородная система имеет ненулевое решение, так как W обращает в нуль детерминант из коэффициентов

$$\left| \sum_s W^{-l_{ij}^{(s)}} - \delta_{ij} \right| = 0.$$

Определенные таким образом $p_{ij}^{(s)}$ могут служить переходными вероятностями, так как, во-первых,

$$\sum_{j, s} p_{ij}^{(s)} = \sum_{j, s} \frac{B_j}{B_i} W^{-l_{ij}^{(s)}} = \frac{B_j}{B_i} = 1,$$

так что сумма вероятностей выхода из любой фиксированной узловой точки равна 1. Кроме того, они неотрицательны, что можно увидеть при рассмотрении величин A_i , приведенных в приложении 1. Эти A_i всегда неотрицательны, а B_i удовлетворяют аналогичной системе уравнений только с переменной местами i и j . Это приводит к противоположной ориентации линий на графике.

Подставляя эти значения $p_{ij}^{(s)}$ в общее выражение для скорости, получим

$$\begin{aligned} & - \frac{\sum P_i p_{ij}^{(s)} \log \frac{B_j}{B_i} W^{-l_{ij}^{(s)}}}{\sum P_i p_{ij}^{(s)} l_{ij}^{(s)}} = \\ & = \frac{\log W \sum P_i p_{ij}^{(s)} l_{ij}^{(s)} - \sum P_i p_{ij}^{(s)} \log B_j + \sum P_i p_{ij}^{(s)} \log B_i}{\sum P_i p_{ij}^{(s)} l_{ij}^{(s)}} = \\ & = \log W = C. \end{aligned}$$

Следовательно, скорость, достигаемая для этого набора вероятностей перехода, равна C , и так как эта скорость не может быть превышена, то она и является максимальной.

Приложение 5

Пусть S_1 — любое измеримое подмножество ансамбля g , а S_2 — подмножество ансамбля f , которое переходит в S_1 при применении оператора T . Тогда

$$S_1 = TS_2.$$

Обозначим через H^λ оператор, сдвигающий все функции на время λ . Тогда

$$H^\lambda S_1 = H^\lambda TS_2 = TH^\lambda S_2,$$

так как T — инвариантен и, следовательно, перестановочен с H^λ . Таким образом, если $m[S]$ есть вероятностная мера множества S , то

$$m[H^\lambda S_1] = m[TH^\lambda S_2] = m[H^\lambda S_2] = m[S_2] = m[S_1],$$

где второе равенство следует из определения меры в пространстве g , третье — из стационарности ансамбля f , а последнее опять-таки

из определения меры g . Это показывает, что ансамбль g стационарен.

Для доказательства сохранения эргодического свойства при применении инвариантных операторов обозначим через S_1 подмножество ансамбля g , инвариантное относительно H^λ , а через S_2 — полный прообраз S_1 для оператора T . Тогда

$$H^\lambda S_1 = H^\lambda T S_2 = T H^\lambda S_2 = S_1,$$

так что $H^\lambda S_1$ включается в S_2 при всех λ . Теперь, так как

$$m[H^\lambda S_2] = m[S_2] = m[S_1],$$

то

$$H^\lambda S_2 = S_2$$

для всех λ при $m[S_2] \neq 0, 1$. Это противоречие показывает, что такого S_1 не существует.

Приложение 6

Верхняя граница, $\bar{N}_3 \leq N_1 + N_2$ получается из-за того, что максимальная возможная энтропия при мощности $N_1 + N_2$ достигается для белого шума указанной мощности. В этом случае энтропийная мощность равна $N_1 + N_2$.

Чтобы получить нижнюю границу, предположим, что имеются два n -мерных распределения $p(x_i)$ и $q(x_i)$ с энтропийными мощностями $\bar{N}_1$ и $\bar{N}_2$. Каковы должны быть $p(x_i)$ и $q(x_i)$ для того, чтобы минимизировать энтропийную мощность $\bar{N}_3$ их свертки $r(x_i)$:

$$r(x_i) = \int p(y_i) q(x_i - y_i) dy_i?$$

Энтропия H_3 распределения r равна

$$H_3 = - \int r(x_i) \log r(x_i) dx_i.$$

Требуется минимизировать ее при условиях

$$H_1 = - \int p(x_i) \log p(x_i) dx_i,$$

$$H_2 = - \int q(x_i) \log q(x_i) dx_i.$$

Для этого рассмотрим

$$U = - \int [r(x) \log r(x) + \lambda p(x) \log p(x) + \mu q(x) \log q(x)] dx;$$

$$\delta U = - \int \{ [1 + \log r(x)] \delta r(x) + \lambda [1 + \log p(x)] \delta p(x) + \mu [1 + \log q(x)] \delta q(x) \} dx.$$

Если $p(x)$ варьируется в окрестности точки $x_i = s_i$, то вариация $r(x)$ равна

$$\delta r(x) = q(x_i - s_i)$$

и

$$\delta U = - \int q(x_i - s_i) \log r(x_i) dx_i - \lambda \log p(s_i) = 0,$$

и аналогично для случая, когда варьируется $q(x)$. Поэтому условиями для минимума будут

$$\int q(x_i - s_i) \log r(x_i) dx_i = -\lambda \log p(s_i),$$

$$\int p(x_i - s_i) \log r(x_i) dx_i = -\mu \log q(s_i).$$

Если умножить первое равенство на $p(s_i)$, а второе на $q(s_i)$ и проинтегрировать по s , то получим

$$H_3 = -\lambda H_1,$$

$$H_3 = -\mu H_2$$

или, решая относительно λ и μ и подставляя в уравнения, получаем, что

$$H_1 \int q(x_i - s_i) \log r(x_i) dx_i = -H_3 \log p(s_i),$$

$$H_2 \int p(x_i - s_i) \log r(x_i) dx_i = -H_3 \log q(s_i).$$

Допустим теперь, что $p(x_i)$ и $q(x_i)$ — нормальные распределения

$$p(x_i) = \frac{|A_{ij}|^{\frac{n}{2}}}{(2\pi)^{\frac{n}{2}}} \exp \left\{ -\frac{1}{2} \sum A_{ij} x_i x_j \right\},$$

$$q(x_i) = \frac{|B_{ij}|^{\frac{n}{2}}}{(2\pi)^{\frac{n}{2}}} \exp \left\{ -\frac{1}{2} \sum B_{ij} x_i x_j \right\}.$$

Тогда $r(x_i)$ также будет нормальным распределением с квадратичной формой c_{ij} . Если обратные формы для A_{ij} , B_{ij} и C_{ij} обозначить через a_{ij} , b_{ij} , c_{ij} , то

$$c_{ij} = a_{ij} + b_{ij}.$$

Требуется показать, что условия минимизации выполняются тогда и только тогда, когда $a_{ij} = Kb_{ij}$, и при этом достигается минимум величины H_3 при наложенных ограничениях. Прежде всего имеем

$$\log r(x_i) = \frac{n}{2} \log \frac{1}{2\pi} |C_{ij}| - \frac{1}{2} \sum C_{ij} x_i x_j,$$

$$\int q(x_i - s_i) \log r(x_i) dx_i = \frac{n}{2} \log \frac{1}{2\pi} |C_{ij}| - \frac{1}{2} \sum C_{ij} s_i s_j - \frac{1}{2} \sum C_{ij} b_{ij},$$

что должно равняться $\frac{H_3}{H_1} \left[\frac{n}{2} \log \frac{1}{2\pi} |A_{ij}| - \frac{1}{2} \sum A_{ij} s_i s_j \right]$, для чего требуется

$$A_{ij} = \frac{H_1}{H_3} C_{ij}.$$

В этом случае $A_{ij} = H_1/H_2 B_{ij}$ и оба уравнения сводятся к тождествам.

Приложение 7

В этом приложении укажем более общий и более строгий подход к основным определениям теории связи¹⁾. Рассмотрим пространство с вероятностной мерой, элементы которого являются упорядоченными парами (x, y) . Величины x, y должны быть отождествлены с возможными передаваемыми и принимаемыми сигналами большой длительности T . Назовем множество всех точек x , для которых x принадлежит к подмножеству S_1 , полосой над S_1 , и аналогично множество, для которого y принадлежит к S_2 , — полосой над S_2 . Разделим x и y на совокупность неперекрывающихся измеримых подмножеств X_i и Y_i , аппроксимируя скорость передачи R выражением

$$R_1 = \frac{1}{T} \sum_i P(X_i, Y_i) \log \frac{P(X_i, Y_i)}{P(X_i)P(Y_i)},$$

где

$P(X_i)$ — вероятностная мера полосы над X_i ,

$P(Y_i)$ — вероятностная мера полосы над Y_i ,

$P(X_i, Y_i)$ — вероятностная мера пересечения этих полос. Дальнейшие подразделения не могут уменьшить R_1 . В самом деле,

¹⁾ Подробное изложение развиваемых здесь идей можно найти в гл. II работы Р. Л. Добрушина, см. сноску на стр. 281. — *Прим. ред.*

пусть X_1 разделяется на два подмножества $X_1 = X'_1 + X''_1$ и

$$P(Y_1) = a, \quad P(X_1) = b + c,$$

$$P(X'_1) = b, \quad P(X'_1, Y_1) = d,$$

$$P(X''_1) = c, \quad P(X''_1, Y_1) = e,$$

$$P(X_1, Y_1) = d + e.$$

Тогда в сумме R_1 заменим (для пересечения X_1 с Y_1)

$$(d + e) \log \frac{d + e}{a(b + c)}$$

на

$$d \log \frac{d}{ab} + e \log \frac{e}{ac}.$$

Легко показать, что при тех ограничениях, которые наложены на b, c, d, e ,

$$\left[\frac{d + e}{b + c} \right]^{d+e} \leq \frac{d^d e^e}{b^d c^e}$$

и, следовательно, сумма возрастет. Таким образом, различные возможные разбиения образуют упорядоченное множество и при этом R_1 монотонно возрастает с утончением разбиения. Можно однозначно определить R как наименьшую верхнюю границу для R_1 и записать ее в виде

$$R = \frac{1}{T} \iint P(x, y) \log \frac{P(x, y)}{P(x)P(y)} dx dy.$$

Этот интеграл, понимаемый в описанном выше смысле, включает как случай непрерывной, так и дискретной систем, а также и многие другие случаи, которые не могут быть представлены ни в одной из этих форм. Из такой формулировки тривиально следует, что если x и u находятся во взаимнооднозначном соответствии, то скорость при передаче от u к y равна скорости передачи от x к y . Если v — любая функция от y (не обязательно имеющая обратную), то скорость передачи от x к v больше или равна скорости передачи от x к y , так как при вычислении приближений любое разбиение для y тоньше аналогичного разбиения для v . В более общем случае, если y и v связаны не функционально, а статистически, т. е. в пространстве (y, v) имеется вероятностная мера, то $R(x, v) \leq R(x, y)$. Это означает, что любая операция, примененная к принимаемому сигналу, даже включающая статистические элементы, не увеличивает R .

Другое понятие, которое должно быть точно определено при абстрактной формулировке теории, — это понятие «скорости числа измерений», которая представляет собой среднее число измерений, требуемых в секунду для того, чтобы задать некоторый член

ансамбля¹⁾. В том случае когда полоса ограничена, достаточно $2W$ чисел в секунду. Общее определение может быть построено следующим образом. Пусть $f_\alpha(t)$ — ансамбль функций и $q_T[f_\alpha(t), f_\beta(t)]$ — метрическое измерение «расстояния» от $f_\alpha(t)$ до $f_\beta(t)$ за время T (например, средне-квадратичное отклонение на этом интервале). Пусть $N(\epsilon, \delta, T)$ — наименьшее число элементов f , которые могут быть выбраны таким образом, что каждый элемент этого ансамбля, за исключением множества меры δ , находится на расстоянии, меньшем ϵ , по крайней мере от одного из выбранных элементов. Таким образом, все пространство покрывается с точностью до ϵ , за исключением множества малой меры δ . Определим скорость числа измерений λ для этого ансамбля с помощью тройного предела:

$$\lambda = \lim_{\delta \rightarrow 0} \lim_{\epsilon \rightarrow 0} \lim_{T \rightarrow 0} \frac{\log N(\epsilon, \delta, T)}{T \log \epsilon}.$$

Это определение является обобщением, использующим теорию меры для определения числа измерений в топологии, и согласуется с интуитивно представляемой скоростью числа измерений для простых ансамблей, где желаемый результат очевиден.

¹⁾ См. статью: Колмогоров А. Н. и Тихомиров В., ϵ -энтропия и ϵ -емкость множеств в метрических пространствах, *Успехи матем. наук*, 14 (1959) 2, 3.—Прим. ред.

Материал, изложенный в данной статье, первоначально составлял содержание секретного доклада «Математическая теория криптографии», датированного 1 сентября 1945 г., который в настоящее время рассекречен.

1. Введение и краткое содержание

Вопросы криптографии и секретных систем открывают возможность для интересных применений теории связи²⁾. В настоящей статье развивается теория секретных систем. Изложение ведется в теоретическом плане и имеет своей целью дополнить положения, приводимые в обычных работах по криптографии³⁾. В этих работах детально изучаются многие стандартные типы кодов и шифров, а также способы их расшифровки. Мы будем иметь дело с общей математической структурой и свойствами секретных систем.

Наше изложение будет ограничено в нескольких отношениях. Во-первых, имеются три общих типа секретных систем: 1) системы маскировки, которые включают применение таких методов, как невидимые чернила, представление сообщения в форме безобидного текста или маскировки криптограммы, и другие методы, при помощи которых факт наличия сообщения скрывается от противника; 2) тайные системы (например, инвертирование речи), в которых для раскрытия сообщения требуется специальное оборудование; 3) «собственно» секретные системы, где смысл сообщения скрывается при помощи шифра, кода и т. д., но само существование сообщения не скрывается и предполагается, что противник обладает любым специальным оборудованием, необходимым для перехвата

¹⁾ Shannon C., Communication theory of secrecy systems, *Bell System Techn. J.*, 28, № 4 (1949), 656—715.

²⁾ См. стр. 243 данного сборника.

³⁾ См., например, Gaines H. F., *Elementary cryptanalysis*, или Giliege M., *Cours de cryptographie*.

и записи переданных сигналов. Здесь будет рассмотрен только третий тип систем, так как системы маскировки представляют в основном психологическую проблему, а тайные системы — техническую проблему.

Во-вторых, наше изложение будет ограничено случаем дискретной информации, где сообщение, которое должно быть зашифровано, состоит из последовательных дискретных символов, каждый из которых выбран из некоторого конечного множества. Эти символы могут быть буквами или словами некоторого языка, амплитудными уровнями «квантованной» речи или видеосигнала и т. д., но главное ударение будет сделано на случае букв.

Статья делится на три части. Резюмируем теперь кратко основные результаты исследования. В первой части излагается основная математическая структура секретных систем. В теории связи считается, что язык может рассматриваться как некоторый вероятностный процесс, который создает дискретную последовательность символов в соответствии с некоторой системой вероятностей. С каждым языком связан некоторый параметр D , который можно назвать избыточностью этого языка. Избыточность измеряет в некотором смысле, насколько может быть уменьшена длина некоторого текста в данном языке без потери какой-либо части информации. Простой пример: так как в словах английского языка за буквой q всегда следует только буква u , то u может быть без ущерба опущена. Значительные сокращения в английском языке можно осуществить, используя его статистическую структуру, частую повторяемость определенных букв или слов, и т. д. Избыточность играет центральную роль в изучении секретных систем.

Секретная система определяется абстрактно как некоторое множество отображений одного пространства (множества возможных сообщений) в другое пространство (множество возможных криптограмм). Каждое конкретное отображение из этого множества соответствует способу шифрования при помощи конкретного ключа.

Предполагается, что отображения являются взаимнооднозначными, так что если известен ключ, то в результате процесса расшифровки возможен лишь единственный ответ.

Предполагается далее, что каждому ключу (и, следовательно, каждому отображению) соответствует некоторая априорная вероятность — вероятность выбрать этот ключ. Аналогично каждому возможному сообщению соответствует априорная вероятность, определяемая задающим сообщение вероятностным процессом. Эти вероятности различных ключей и сообщений являются фактически априорными вероятностями для шифровальщика противника и характеризуют его априорные знания относительно интересующей его проблемы.

Чтобы использовать такую секретную систему, сначала выбирается некоторый ключ и посылается в точку приема. Выбор ключа определяет конкретное отображение из множества отображений, образующих систему. Затем выбирается сообщение и с помощью отображения, соответствующего выбранному ключу, из этого сообщения формируется криптограмма. Эта криптограмма передается в точку приема по некоторому каналу и может быть перехвачена «противником»¹⁾. На приемном конце с помощью отображения, обратного выбранному, из криптограммы восстанавливают первоначальное сообщение.

Если противник перехватит криптограмму, он может с ее помощью сосчитать апостериорные вероятности различных возможных сообщений и ключей, которые могли быть использованы для составления такой криптограммы. Это множество апостериорных вероятностей образует его сведения о ключах и сообщениях после перехвата. «Сведения», таким образом, представляют собой некоторое множество предположений, которым приписаны вероятности. Вычисление апостериорных вероятностей является общей задачей расшифровки.

Проиллюстрируем эти понятия простым примером. В шифре простой подстановки со случайным ключом имеется $26!$ отображений, соответствующих $26!$ способам, которыми мы можем заменить 26 различных букв²⁾. Все эти способы равновозможны, и поэтому каждый имеет априорную вероятность $1/26!$. Если такой шифр применяется к «нормативному английскому языку» и предполагается, что шифровальщик противника не знает ничего об источнике сообщений, кроме того, что он создает английский текст, то априорными вероятностями различных сообщений из N букв являются просто их относительные частоты в нормативном английском тексте.

Если противник перехватил такую криптограмму из N букв, его апостериорные вероятности изменятся. Если N достаточно велико (скажем, 50 букв), имеется обычно единственное сообщение с апостериорной вероятностью, близкой к единице, в то время как все другие сообщения имеют суммарную вероятность, близкую к нулю. Таким образом, имеется, по существу, единственное «решение» такой криптограммы. Для меньших N (скажем, $N = 15$) обычно найдется много сообщений и ключей, вероятности которых сравнимы, и не найдется ни одного сообщения и ключа с вероятностью, близкой к единице. В этом случае «решение» криптограммы неоднозначно.

¹⁾ Слово «противник», заимствованное из военных приложений, обычно используется в криптографии для обозначения кого-либо, кто может перехватить криптограмму.

²⁾ Здесь подразумевается шифр, в котором каждая из 26 букв латинского алфавита заменяется другой буквой того же алфавита. — *Прим. ред.*

В результате рассмотрения секретных систем, которые могут быть представлены как совокупность отображений одного множества элементов в другое, возникают две естественные операции комбинирования, производящие из двух данных систем третью. Первая операция комбинирования называется операцией «умножения» (произведением) и соответствует зашифровке сообщения с помощью системы R с последующей зашифровкой полученной криптограммы с помощью системы S , причем ключи R и S выбираются независимо. Полный результат этой операции представляет собой секретную систему, отображения которой состоят из всех произведений (в обычном смысле произведений отображений) отображений из S на отображения из R . Вероятности результирующих отображений являются произведениями вероятностей двух исходных отображений.

Вторая операция комбинирования является «взвешенным сложением»:

$$T = pR + qS, \quad p + q = 1.$$

Она представляет собой следующее. Сначала делается предварительный выбор, какая из систем R или S будет использоваться, причем система R выбирается с вероятностью p , а система S с вероятностью q . После этого выбранная система используется описанным выше способом.

Будет показано, что секретные системы с этими двумя операциями комбинирования образуют, по существу, «линейную ассоциативную алгебру» с единицей, — алгебраический объект, подробно изучавшийся математиками.

Среди многих возможных секретных систем имеется один тип с многочисленными особыми свойствами. Этот тип назовем «чистой» системой. Система является чистой, если все ключи равновероятны и если для любых трех отображений T_i, T_j, T_k из множества отображений данной системы произведение

$$T_i T_j^{-1} T_k$$

также является отображением из этого множества. То есть зашифровка, расшифровка и снова зашифровка с любыми тремя ключами должна быть эквивалентна зашифровке с некоторым ключом.

Можно показать, что для чистого шифра все ключи по существу эквивалентны — все они приводят к тому же самому множеству апостериорных вероятностей. Больше того, каждой криптограмме соответствует некоторое множество сообщений («остаточный класс»), из которых могла бы получиться эта криптограмма, а апостериорные вероятности сообщений в этом классе пропорциональны априорным вероятностям. Вся информация, которую противник получил бы в результате перехвата криптограммы, заклю-

чается в установлении остаточного класса. Многие из обычных шифров являются чистыми системами, в том числе простая подстановка со случайным ключом. В этом случае остаточный класс состоит из всех сообщений с таким же набором буквенных повторений, как в перехваченной криптограмме.

По определению, две системы R и S являются «подобными», если существует фиксированное отображение A (имеющее обратное A^{-1}), такое, что

$$R = AS.$$

Если R и S подобны, то между получающимися в результате применения этих систем множествами криптограмм можно установить взаимнооднозначное соответствие, приводящее к тем же самым апостериорным вероятностям. Такие две системы аналитически записываются одинаково.

Во второй части статьи рассматривается проблема «теоретической секретности». Насколько легко некоторая система поддается раскрытию при условии, что для анализа перехваченной криптограммы противник располагает неограниченным количеством времени и специалистов? Эта проблема тесно связана с вопросами связи при наличии шумов, и понятия энтропии и неопределенности, введенные в теории связи, находят прямое применение в этом разделе криптографии.

«Совершенная секретность» определяется следующими требованиями к системе. Требуется, чтобы апостериорные вероятности различных сообщений, полученные после перехвата противником данной криптограммы, были бы в точности равны априорным вероятностям тех же сообщений до перехвата. Покажем, что «совершенная секретность» возможна, но требует в случае конечного числа сообщений того же самого числа возможных ключей. Если считать, что сообщение создается с данной «скоростью» R (понятие скорости будет определено позже), то ключ должен создаваться с той же самой или с большей скоростью.

Если используется секретная система с конечным ключом и перехвачены N букв криптограммы, то для противника будет существовать определенное множество сообщений с определенными вероятностями, которые могли бы создать эту криптограмму. С увеличением N это множество обычно сужается до тех пор, пока в конце концов не получится единственного «решения» криптограммы: одно сообщение с вероятностью, близкой к единице, а все остальные с вероятностями, практически равными нулю. В работе определяется величина $H(N)$, названная *ненадежностью*. Эта величина измеряет (в статистическом смысле), насколько близка средняя криптограмма из N букв к единственному решению, т. е. насколько неточно известно противнику истинное сообщение после перехвата

криптограммы из N букв. Далее выводятся различные свойства ненадежности, например: ненадежность ключа не возрастает с ростом N . Эта ненадежность является теоретическим показателем секретности — теоретическим, поскольку она позволяет противнику дешифровать криптограмму лишь в том случае, если он обладает неограниченным запасом времени.

В этой же части определяется функция $H(N)$ для некоторых идеализированных типов шифров, называемых *случайными шифрами*. С некоторыми видоизменениями эта функция может быть применена ко многим случаям, представляющим практический интерес. Это дает способ приближенного вычисления количества материала, который требуется перехватить, чтобы получить решение секретной системы.

Из подобного анализа следует, что для обычных языков и обычных типов шифров (но не кодов) это «расстояние единственности» равно приблизительно $H(K)/D$. Здесь $H(K)$ — число, измеряющее «объем» пространства ключей. Если все ключи априори равновероятны, то $H(K)$ равно логарифму числа возможных ключей. Вводимое число D — это избыточность языка. Оно измеряет количество «статистических ограничений», налагаемых языком. Для простой подстановки со случайным ключом наше $H(K)$ равно $\log_{10} 26!$ или приблизительно 20, а D (в десятичных единицах на букву) для английского языка равно приблизительно 0,7. Таким образом, единственность решения достигается приблизительно при 30 буквах.

Для некоторых «языков» можно построить такие секретные системы с конечным ключом, в которых неопределенность не стремится к нулю при $N \rightarrow \infty$. В этом случае противник не получит единственного решения такого шифра, сколько бы материала он не перехватил, и у него будет оставаться много альтернатив с довольно большими вероятностями. Такие системы назовем *идеальными системами*. В любом языке можно аппроксимировать такую ситуацию, т. е. отсрочить приближение $H(N)$ к нулю до сколь угодно больших N . Однако такие системы имеют много недостатков, таких, как сложность и чувствительность к ошибкам при передаче криптограммы.

Третья часть статьи посвящена «практической секретности». Две системы с одинаковым объемом ключа могут быть обе разрешимы единственным образом, когда перехвачено N букв, но они могут значительно отличаться по количеству времени и усилий, затрачиваемых для получения решения. На основе анализа основных недостатков секретных систем предлагаются методы построения систем, для решения которых требуются большие затраты времени и сил. Наконец, рассматривается проблема несовместимости различных желательных качеств секретных систем.

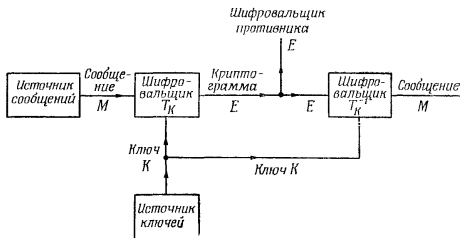
Часть I

МАТЕМАТИЧЕСКАЯ СТРУКТУРА СЕКРЕТНЫХ СИСТЕМ

2. Секретные системы

Чтобы приступить к математическому анализу криптографии, необходимо ввести удовлетворительную идеализацию и определить математически приемлемым способом, что будет пониматься под термином секретная система. Схематическая структура секретной системы показана на рис. 1.

На передающем конце имеются два источника информации — источник сообщений и источник ключей. Источник ключей отбирает



Р и с. 1. Схема общей секретной системы.

конкретный ключ среди всех возможных ключей данной системы. Этот ключ передается некоторым способом на приемный конец, причем предполагается, что его нельзя перехватить (например, ключ передается посыльным). Источник сообщений формирует некоторое сообщение (незашифрованное), которое затем зашифровывается, и готовая криптограмма передается на приемный конец, причем криптограмма может быть перехвачена (например, пересылается по радио). На приемном конце шифровальщик с помощью ключа по криптограмме восстанавливает исходное сообщение.

Очевидно, шифровальщик на передающем конце выполняет некоторую функциональную операцию. Если M — сообщение, K — ключ и E — зашифрованное сообщение (криптограмма), то имеем

$$E = f(M, K),$$

т. е. E является функцией от M и K . Удобнее, однако, понимать E не как функцию двух переменных, а как (однопараметрическое) семейство операций или отображений, и записывать его в виде:

$$E = T_i M.$$

Отображение T_i , примененное к сообщению M , дает криптограмму E . Индекс i соответствует конкретному используемому ключу.

Вообще мы будем предполагать, что имеется лишь конечное число возможных ключей, каждому из которых соответствует вероятность p_i . Таким образом, источник ключей является статистическим процессом, или устройством, которое выбирает одно из множества отображений $T_1, \dots, T_m$ с вероятностями $p_1, \dots, p_m$ соответственно. Будем также предполагать, что число возможных сообщений конечно, и эти сообщения $M_1, \dots, M_n$ имеют априорные вероятности $q_1, \dots, q_n$. Например, возможными сообщениями могли бы быть всевозможные последовательности английских букв, включающих по N букв каждая, а соответствующими вероятностями тогда были бы относительные частоты появления таких последовательностей в нормативном английском тексте.

Должна иметься возможность восстанавливать M на приемном конце, когда известны E и K . Поэтому отображение T_i из нашего семейства должно иметь единственное обратное отображение T_i^{-1} , так что $T_i T_i^{-1} = I$, где I — тождественное отображение. Таким образом:

$$M = T_i^{-1} E.$$

Во всяком случае, это обратное отображение T_i^{-1} должно существовать и быть единственным для каждого E , которое может быть получено из M с помощью ключа i . Приходим, таким образом, к следующему определению: секретная система есть семейство однозначно обратимых отображений T_i множества возможных сообщений во множество криптограмм, при этом отображение T_i имеет вероятность p_i . Обратно, любое множество объектов такого типа будет называться «секретной системой». Множество возможных сообщений для удобства будет называться «пространством сообщений», а множество возможных криптограмм — «пространством криптограмм».

Две секретные системы совпадают, если они образованы одним и тем же множеством отображений T_i и одинаковыми пространствами сообщений и криптограмм, причем вероятности ключей в этих системах также совпадают.

Секретную систему можно представлять себе как некоторую машину с одним или более переключающими устройствами. Последовательность букв (сообщение) поступает на вход машины, а на выходе ее получается другая последовательность. Конкретное

положение переключающих устройств соответствует конкретному используемому ключу. Для выбора ключа из множества возможных ключей должны быть заданы некоторые статистические методы.

Для того чтобы нашу проблему можно было рассмотреть математически, предположим, что противнику известна используемая система. Иными словами, он знает семейство отображений T_i и вероятности выбора различных ключей. Можно было бы, во-первых, возразить, что такое предположение нереалистично, так как шифровальщик противника часто не знает, какая система использовалась или чему равны рассматриваемые вероятности. На это возражение имеется два ответа.

1. Наложённое ограничение слабее, чем кажется с первого взгляда, из-за широты нашего определения секретной системы. Предположим, что шифровальщик перехватывает сообщение и не знает, использовалась ли здесь подстановка, или транспозиция, или шифр типа Виженера. Он может считать, что сообщение зашифровано с помощью системы, в которой часть ключа является указанием того, какой из трех типов имеющихся ключей был использован, а следующая часть — конкретный ключ этого типа. Указанным трем различным возможностям шифровальщик приписывает вероятности, учитывая при этом все имеющиеся у него сведения об априорных вероятностях использования шифровальщиком противника соответствующих типов шифров.

2. Наше ограничение обычно в криптографических исследованиях. Оно является пессимистичным, но безопасно, и в конечном счете реалистично, так как можно ожидать, что противник рано или поздно раскроет любую секретную систему. Поэтому даже в том случае, когда разработана совершенно новая система, так что противник не может приписать ей никаких априорных вероятностей, если только он ее уже не раскрыл, нужно иметь в виду его возможную осведомленность.

Эта ситуация аналогична ситуации, возникающей в теории игр¹⁾, где предполагается, что партнер «обнаруживает» используемую стратегию игры. В обоих случаях это предположение служит для более четкого описания сведений, которыми располагает противная сторона.

Второе возможное возражение против нашего определения секретной системы состоит в том, что в нем не принимаются в расчет используемые обычно на практике вставки в сообщение посторон-

¹⁾ См. von Neumann J., Morgenstern O., The theory of games and economical behavior, Princeton, 1947. [Из современной литературы см., например, книгу Ма к-К и н с и Дж., Введение в теорию игр, Физматгиз, 1960 или Л ь ю с Д., Р а й ф ф а Х., Игры и решения, ИЛ, 1961.—Прим. ред.]

них нулевых знаков и использование многократных подстановок. В таких случаях для данного сообщения и ключа имеется не единственная криптограмма, и шифровальщик может выбрать по своему желанию одну из нескольких различных криптограмм. Эту ситуацию можно было бы рассмотреть, но это только внесло бы дополнительные усложнения на данном этапе рассуждений без существенного изменения каких-либо из основных выводов.

Если сообщения создаются марковским процессом [типа, описанного в работе, указанной в примечании на стр. 341], то вероятности разных сообщений определяются структурой этого марковского процесса. Однако подойдем к вопросу с более общей точки зрения и будем трактовать сообщения просто как абстрактное множество объектов, которым приписаны вероятности, причем эти объекты не обязательно состоят из последовательностей букв и не обязательно создаются марковским процессом.

Следует подчеркнуть, что далее во всех случаях секретная система означает не одно, а целое множество отображений. После того как выбран ключ, используется только одно из этих отображений и отсюда можно было бы прийти к определению секретной системы как единственного преобразования языка. Однако противник не знает, какой ключ выбран и остальные возможные ключи столь же важны для него, как и истинный. Именно существование этих других возможных ключей и придает системе секретность. Так как мы интересуемся в первую очередь секретностью, то вынуждены предпочесть данное нами определение понятия секретной системы. Тип ситуации, когда остальные возможности так же важны, как и осуществившаяся, часто встречается в стратегических играх. Ход шахматной игры в большой степени контролируется угрозами, которые не осуществляются. Нечто подобное представляет из себя «фактическое существование» нереализованных возможностей в теории игр.

Следует отметить, что система, состоящая из единственной операции над языком, представляет собой при нашем определении вырожденный тип секретной системы. Это — система с единственным ключом, который имеет вероятность, равную единице. В такой системе нет секретности — шифровальщик противника находит сообщение, применяя к перехваченной криптограмме обратное отображение, также единственное в такой системе. В этом случае шифровальщик противника и шифровальщик получателя информации располагают одинаковой информацией. В общем же случае единственное различие их сведений состоит в том, что последнему известен конкретно использовавшийся ключ, в то время как первому известны лишь априорные вероятности различных ключей из данного множества. Процесс расшифровки для получателя информации состоит в применении к криптограмме отображения, обратного по отноше-

нию к конкретному отображению, использованному для составления криптограммы. Процесс расшифровки для противника представляет собой попытку определить сообщение (или конкретный ключ), имея в распоряжении только криптограмму и априорные вероятности различных ключей и сообщений.

Существует много трудных эпистемологических вопросов, связанных с теорией секретности, или вернее с любой теорией, связанной с реальным применением вопросов теории вероятностей (так обстоит дело, в частности, с априорными вероятностями, теоремой Байеса и т. д.). Трактующая абстрактно теория вероятности может быть изложена на строгих логических основах с использованием современной теории меры¹⁾. Однако в приложениях к физическим ситуациям, особенно когда дело касается «субъективных» вероятностей и неповторимых экспериментов, возникают многочисленные вопросы, связанные с логическим обоснованием. Например, при нашем подходе к проблеме секретности допускается, что априорные вероятности различных ключей и сообщений известны шифровальщику противника, но как он может определить их эффективным способом даже при использовании всех своих сведений о данной обстановке?

Можно создать искусственные криптографические ситуации типа «урны и игральной кости», в которых априорные вероятности имеют вполне определенный смысл, и идеализация, использованная здесь, является наверняка подходящей. Но в других случаях, которые можно себе представить, например, при перехвате сообщений, передаваемых между собой марсианами, высадившимися на землю, априорные вероятности были бы настолько неопределенными, что не имели бы никакого значения.

Наиболее часто встречающиеся на практике криптографические задачи лежат где-то между этими крайними пределами. Шифровальщик противника может иметь желание разделить возможные сообщения на категории «приемлемых», «возможных, но маловероятных» и «неприемлемых», но чувствуется, что более подробное подразделение не имело бы смысла.

К счастью, на практике только очень большие ошибки в априорных вероятностях ключей и сообщений могут вызвать заметные ошибки в важных параметрах. Это происходит из-за того, что число сообщений и криптограмм ведет себя как экспоненциальная функция, а измеряется логарифмической мерой.

¹⁾ См., D o o b J. L., Probability as measure, *Ann. of Math. Stat.*, 12 (1941), 206—214 [см. также Д у б Дж. Л., Вероятностные процессы, ИЛ, М., 1956.—Прим. ред.]; K o l m o g o r o f f A., Grundbegriffe der Wahrscheinlichkeitsrechnung, *Ergebnisse der Mathematik*, v. 2, № 3, Berlin, 1933. [К о л м о г о р о в А. Н., Основные понятия теории вероятностей, М., ОНТИ, 1936.—Прим. ред.]

3. Способы изображения систем

Секретная система, в том виде как она определена выше, может быть изображена различными способами. Один из них (удобный для целей иллюстрации) использует линейные схемы, изображенные на рис. 2 и рис. 4. Возможные сообщения представляются точками слева, а возможные криптограммы — точками справа. Если некоторый ключ, скажем ключ 1 , отображает сообщение M_2 в криптограмму E_4 , то M_2 и E_4 соединяются линией, обозначенной значком 1 и т. д. Для каждого ключа из каждого сообщения должна

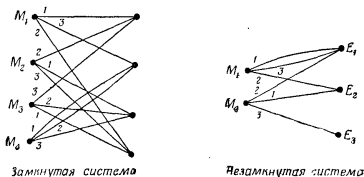


Рис. 2. Схемы простых систем

выходить ровно одна линия. Если это же верно и для каждой криптограммы, скажем, что система является *замкнутой*.

Более общий способ описания системы состоит в задании операции, с помощью которой, применяя к сообщению произвольный ключ, можно получить криптограмму. Аналогично неявным образом можно определить вероятности различных ключей или с помощью задания способа выбора ключей, или с помощью описания сведений о том, как обычно выбирает ключи противник. Вероятности сообщений определяются просто посредством изложения наших априорных сведений о языке противника, тактической обстановке (которая будет влиять на возможное содержание сообщений) и любой специальной информации, касающейся криптограммы.

4. Примеры секретных систем

В данном разделе рассматриваются несколько примеров шифров. В дальнейшем в целях иллюстрации будем часто ссылаться на эти примеры.

1. Шифр простой подстановки.

В таком шифре производится замена каждой буквой сообщения на некоторый определенный символ (обычно также на букву).

Таким образом, сообщение

$$M = m_1 m_2 m_3 m_4 \dots,$$

где $m_1, m_2, \dots$ — последовательные буквы, переходит в

$$E = e_1 e_2 e_3 e_4 \dots = f(m_1) f(m_2) f(m_3) f(m_4) \dots,$$

причем функция $f(m)$ имеет обратную функцию. Ключ является просто перестановкой алфавита (если буквы заменяются на буквы), например, *XGUACDTBFHRSMLMQVYZWIEJOKNP*. Первая буква — X заменяет букву A, G заменяет B и т. д.

2. Транспозиция с фиксированным периодом d .

В этом случае сообщение делится на группы символов длины d , и к каждой группе применяется одна и та же перестановка. Эта перестановка является ключом; она может быть задана некоторой перестановкой первых d целых чисел. Таким образом, для $d = 5$, в качестве перестановки можно взять 2 3 1 5 4. Это будет означать, что

$$m_1 m_2 m_3 m_4 m_5 m_6 m_7 m_8 m_9 m_{10} \dots$$

переходит в

$$m_2 m_3 m_1 m_5 m_4 m_7 m_8 m_6 m_{10} m_9 \dots$$

Последовательное применение двух или более транспозиций будет называться составной транспозицией. Если периоды этих транспозиций $d_1, \dots, d_s$, то, очевидно, в результате получится транспозиция периода d , где d — наименьшее общее кратное $d_1, \dots, d_s$.

3. Шифр Виженера и его варианты.

В шифре Виженера ключ задается набором из d букв. Такие наборы подписываются с повторением под сообщением и полученные две последовательности складываются по модулю 26 (каждая буква рассматриваемого алфавита нумеруется от $A = 0$ до $Z = 25$). Таким образом,

$$l_i = m_i + k_i \pmod{26},$$

где k_i — буква ключа, полученная сокращением числа i по модулю d . Например, с помощью ключа *GAH* получаем

Сообщение		N	O	W	I	S	T	H	E
Повторяемый ключ		G	A	H	G	A	H	G	A
Криптограмма		T	O	D	O	S	A	N	E

Шифр Виженера с периодом 1 называется шифром Цезаря. Он представляет собой простую подстановку, в которой каждая буква сообщения M сдвигается вперед на фиксированное число мест по алфавиту. Это число и является ключом; оно может быть любым от 0 до 25. Так называемый шифр Бофора (Beaufort) и видоизме-

ненный шифр Бофора подобны шифру Виженера. В них сообщения зашифровываются с помощью равенств

$$l_i = k_i - m_i \pmod{26}$$

и

$$l_i = m_i - k_i \pmod{26}$$

соответственно. Шифр Бофора с периодом 1 называется обратным шифром Цезаря.

Повторное применение двух или более шифров Виженера будет называться составным шифром Виженера. Он имеет уравнение

$$l_i = m_i + k_i + l_i + \dots + s_i \pmod{26},$$

где $k_i, l_i, \dots, s_i$, вообще говоря, имеют различные периоды. Период их суммы

$$k_i + l_i + \dots + s_i,$$

как и в составной транспозиции, будет наименьшим общим кратным отдельных периодов.

Если используется шифр Виженера с неограниченным неповторяющимся ключом, то мы имеем шифр Вернама¹⁾, в котором

$$l_i = m_i + k_i \pmod{26}$$

и k_i выбираются случайно и независимо среди чисел $0, 1, \dots, 25$. Если ключом служит текст, имеющий смысл, то имеем шифр «бегущего ключа».

4. Диграммная, триграммная и n -граммная подстановки.

Вместо подстановки одной буквы можно использовать подстановку диграмм, триграмм и т. д. Для диграммной подстановки в общем виде требуется ключ, состоящий из перестановок 26^2 диграмм. Он может быть представлен с помощью таблицы, в которой ряд соответствует первой букве диграммы, а столбец — второй букве, причем клетки таблицы заполнены заменяющими символами (обычно также диграммами).

5. Шифр Виженера с перемешанным один раз алфавитом.

Такой шифр представляет собой простую подстановку с последующим применением шифра Виженера

$$l_i = f(m_i) + k_i,$$

$$m_i = f^{-1}(l_i - k_i).$$

¹⁾ V e r n a m G. S., Cipher printing telegraph systems for secret wire and radio telegraphic communication, *J. Am. Inst. Electr. Eng.*, 45 (1926), 109—115.

«Обратным» к такому шифру является шифр Виженера с последующей простой подстановкой

$$l_i = g(m_i + k_i),$$

$$m_i = g^{-1}(l_i) - k_i.$$

6. Матричная система¹⁾

Имеется один метод подстановки n -грамм, который заключается в применении к последовательным n -граммам некоторой матрицы, имеющей обратную. Предполагается, что буквы занумерованы от 0 до 25 и рассматриваются как элементы некоторого алгебраического кольца. Если к n -грамме сообщения применить матрицу a_{ij} , то получится n -грамма криптограммы

$$l_i = \sum_{j=1}^n a_{ij} m_j \quad i = 1, \dots, n.$$

Матрица a_{ij} является ключом и расшифровка выполняется с помощью обратной матрицы. Обратная матрица будет существовать тогда и только тогда, когда определитель $|a_{ij}|$ имеет обратный элемент в нашем кольце.

7. Шифр Плэйфер (Playfair)²⁾

Этот шифр является частным видом диграммной подстановки, которая производится с помощью перемешанного алфавита из 25 букв, записанных в виде квадрата 5×5 . (Буква J часто опускается при криптографической работе, так как она редко встречается, и в тех случаях, когда она встречается, ее можно заменить буквой I). Предположим, что ключевой квадрат записывается следующим образом:

L	Z	Q	C	P
A	G	N	O	U
R	D	M	I	F
K	Y	H	V	S
X	B	T	E	W

В этом случае диграмма AC , например, заменяется на пару букв, расположенных в противоположных углах прямоугольника, определяемого буквами A и C , т. е. на LO , причем L взята первой, так как она выше A . Если буквы диграммы расположены на одной горизонтали, как, например, RI , то используются стоящие справа от них буквы DF ; RF заменяется на DR . Если буквы расположены

¹⁾ См. Hill L. S., Cryptography in an algebraic alphabet, *Amer. Math. Monthly*, 36 (1929), № 6, 1, 306—312, а также «Concerning certain linear transformation apparatus of cryptography», там же, 38 (1931), № 3, 135—154.

²⁾ Предложен Виттстоном. — *Прим. ред.*

на одной вертикали, то используются буквы, стоящие под ними. Таким образом, PS заменяется на UW . Если обе буквы диграммы совпадают, то можно использовать для разделения их нуль или же одну из букв опустить и т. п.

8. *Перемешивание алфавита с помощью многократной подстановки.*

В этом шифре используются последовательно d простых подстановок. Так, если $d = 4$, то

$$m_1 m_2 m_3 m_4 m_5 m_6 \dots$$

заменяется на

$$f_1(m_1) f_2(m_2) f_3(m_3) f_4(m_4) f_1(m_5) f_2(m_6) \dots$$

и т. д.

9. *Шифр с автоключом.*

Шифр типа Виженера, в котором или само сообщение или результирующая криптограмма используются в качестве «ключа», называется шифром с автоключом. Шифрование начинается с помощью «первичного ключа» (который является настоящим ключом в нашем смысле) и продолжается с помощью сообщения или криптограммы, смещенной на длину первичного ключа, как в указанном ниже примере, где первичным ключом является набор букв *COMET*. В качестве «ключа» используется сообщение:

Сообщение	$S E N D S U P P L I E S \dots$
Ключ	$C O M E T S E N D S U P \dots$
Криптограмма	$U S Z H L M T C O A Y H \dots$

Если в качестве «ключа» использовать криптограмму, то получится¹⁾

Сообщение	$S E N D S U P P L I E S \dots$
Ключ	$C O M E T U S Z H L O H \dots$
Криптограмма	$U S Z H L O H O S T T S \dots$

10. Дробные шифры

В этих шифрах каждая буква сначала зашифровывается в две (или более) буквы или в два (или более) числа, затем полученные символы каким-либо способом перемешиваются (например, с помощью транспозиции), после чего их можно снова перевести в первоначальный алфавит. Таким образом, используя в качестве ключа перемешанный 25-буквенный алфавит, можно перевести буквы в двух-

¹⁾ Эта система является тривиальной с точки зрения секретности, так как, за исключением первых d букв, в распоряжении противника имеется весь «ключ».

значные пятеричные числа с помощью таблицы:

	0	1	2	3	4
0	L	Z	Q	C	P
1	A	G	N	O	U
2	R	D	M	I	F
3	K	Y	H	V	S
4	X	B	T	E	W

Например, букве *B* соответствует «число» 41. После того как полученный ряд чисел подвергнут некоторой перестановке, его можно снова разбить на пары чисел и перейти к буквам.

11. Коды.

В кодах слова (или иногда слоги) заменяются группами букв. Иногда затем применяется шифр того или иного вида.

5. Оценка секретных систем

Имеется несколько различных критериев, которые можно было бы использовать для оценки качества предлагаемой секретной системы. Рассмотрим наиболее важные из этих критериев.

1. Количество секретности.

Некоторые секретные системы являются совершенными в том смысле, что положение противника не облегчается в результате перехвата любого количества сообщений. Другие системы, хотя и дают противнику некоторую информацию при перехвате очередной криптограммы, но не допускают единственного «решения». Системы, допускающие единственное решение, очень разнообразны как по затрате времени и сил, необходимых для получения этого решения, так и по количеству материала, который необходимо перехватить для получения единственного решения.

2. Объем ключа.

Ключ должен быть передан из передающего пункта в приемный пункт таким способом, чтобы его нельзя было перехватить. Иногда его нужно запомнить. Поэтому желательно иметь ключ настолько малый, насколько это возможно.

3. Сложность операции шифрования и дешифрирования.

Операции шифрования и дешифрирования должны быть, конечно, по возможности простыми. Если эти операции производятся вручную, то их сложность приводит к потере времени, появлению ошибок и т. д. Если они производятся механически, то сложность приводит к использованию больших и дорогих устройств.

4. Разрастание числа ошибок.

В некоторых типах шифров ошибка в одной букве, допущенная при шифровании или передаче, приводит к большому числу ошибок

в расшифрованном тексте. Такие ошибки разрастаются в результате операции дешифрирования, вызывая значительную потерю информации и часто требуя повторной передачи криптограммы. Естественно, желательно минимизировать это возрастание числа ошибок.

5. Увеличение объема сообщения.

В некоторых типах секретных систем объем сообщения увеличивается в результате операции шифрования. Этот нежелательный эффект можно наблюдать в системах, в которых делается попытка попутить статистику сообщения в массе добавляемых нулевых символов, или где используются многократные замены. Он имеет место также во многих системах типа «маскировки» (которые не являются обычными секретными системами в смысле нашего определения).

6. Алгебра секретных систем

Если имеются две секретные системы T и R , их часто можно комбинировать различными способами для получения новой секретной системы S . Если T и R имеют одну и ту же область (пространство сообщений), то можно образовать своего рода «взвешенную сумму»

$$S = pT + qR,$$

где $p + q = 1$. Эта операция состоит, во-первых, из предварительного выбора систем T или R с вероятностями p и q . Этот выбор является частью ключа S . После того как этот выбор сделан, системы T или R применяются в соответствии с их определениями. Полный ключ S должен указывать, какая из систем T или R выбрана и с каким ключом используется выбранная система.

Если T состоит из отображения $T_1, \dots, T_m$ с вероятностями $p_1, \dots, p_m$, а R из $R_1 \dots R_k$ с вероятностями $q_1, \dots, q_k$, то система $S = pT + qR$ состоит из отображений $T_1, \dots, T_m, R_1, \dots, R_k$ с вероятностями $pp_1, \dots, pp_m, qq_1, \dots, qq_k$ соответственно.

Обобщая далее, можно образовать сумму нескольких систем

$$S = p_1T + p_2R + \dots + p_mU, \quad \sum p_i = 1.$$

Заметим, что любая система T может быть записана как сумма фиксированных операций

$$T = p_1T_1 + p_2T_2 + \dots + p_mT_m,$$

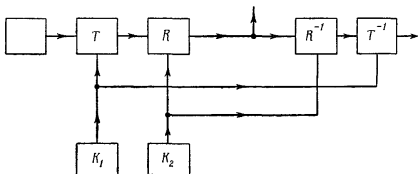
где T_i — определенная операция шифрования в системе T , соответствующая выбору ключа i , причем вероятность такого выбора равна p_i .

Второй способ комбинирования двух секретных систем заключается в образовании «произведения», как показано схематически

на рис. 3. Предположим, что T и R — такие две системы, что область определения (пространство языка) системы R может быть отождествлена с областью определения (пространством криптограмм) системы T . Тогда можно применить сначала систему T к нашему языку, а затем систему R к результату этой операции, что дает результирующую операцию S , которую запишем в виде произведения

$$S = RT.$$

Ключ системы S состоит как из ключа системы T , так и из ключа системы R , причем предполагается, что эти ключи выбираются соответственно их первоначальным вероятностям и независимо.



Р и с. 3. Произведение двух систем $S = RT$.

Таким образом, если m ключей системы T выбирается с вероятностями

$$p_1, p_2, \dots, p_m,$$

а n ключей системы R имеют вероятности

$$p'_1, p'_2, \dots, p'_n,$$

то система S имеет самое большее mn ключей с вероятностями $p_i p'_j$. Во многих случаях некоторые из отображений $R_i T_j$ будут одинаковыми и могут быть сгруппированы вместе, а их вероятности при этом сложатся.

Произведение шифров используется часто; например, после подстановки применяют транспозицию или после транспозиции — код Виженера; или же применяют код к тексту и зашифровывают результат с помощью подстановки, транспозиции, дробным шифром и т. д.

Можно заметить, что такое умножение, вообще говоря, некоммутативно (т. е. не всегда $RS = SR$), хотя в частных случаях (таких, как подстановка и транспозиция) коммутативность имеет место.

Так как наше умножение представляет собой некоторую операцию, оно по определению ассоциативно, т. е. $R(ST) = (RS)T = = RST$. Кроме того, верны законы

$$p(p'T + q'R) + qS = pp'T + pq'R + qS$$

(взвешенный ассоциативный закон для сложения);

$$T(pR + qS) = pTR + qTS$$

$$(pR + qS)T = pRT + qST$$

(право- и левосторонние дистрибутивные законы), а также справедливо равенство

$$p_1T + p_2T + p_3R = (p_1 + p_2)T + p_3R.$$

Следует подчеркнуть, что эти операции комбинирования сложения и умножения применяются к секретным системам в целом. Произведение двух систем TR не следует смешивать с произведением отображений в системах T_iR_j , которое также часто используется в настоящей работе. Первое является секретной системой, т. е. множеством отображений с соответствующими вероятностями; второе — является фиксированным отображением. Далее, в то время как сумма двух систем $pR + qT$ является системой, сумма двух отображений не определена. Системы T и R могут коммутировать, в то время как конкретные R_j и T_i не коммутируют. Например, если R — система Бофора данного периода, все ключи которой равновероятны, то, вообще говоря,

$$R_iR_j \neq R_jR_i,$$

но, конечно, произведение RR не зависит от порядка сомножителей; действительно

$$RR = V$$

является системой Виженера того же самого периода со случайным ключом. С другой стороны, если отдельные отображения T_i и R_j двух систем T и R коммутируют, то и системы коммутируют.

Системы, у которых пространства M и E можно отождествить (этот случай является очень частым, если последовательности букв преобразуются в последовательности букв), могут быть названы *эндоморфными*. Эндоморфная система T может быть возведена в степень T^n .

Секретная система T , произведение которой на саму себя равно T , т. е. такая, что

$$TT = T,$$

будет называться *идемпотентной*. Например, простая подстановка, транспозиция с периодом p , система Виженера с периодом p (все с равновероятными ключами) являются идемпотентными.

Множество всех эндоморфных секретных систем, определенных в фиксированном пространстве сообщений, образует «алгебраическую систему», т. е. некоторый вид алгебры, использующий операции сложения и умножения. Действительно, рассмотренные свойства сложения и умножения можно резюмировать следующим образом.

Множество эндоморфных шифров с одним и тем же пространством сообщений и двумя операциями комбинирования — операцией взвешенного сложения и операцией умножения — образуют линейную ассоциативную алгебру с единицей, с той лишь особенностью, что коэффициенты во взвешенном сложении должны быть неотрицательными, а их сумма должна равняться единице.

Эти операции комбинирования дают способы конструирования многих новых типов секретных систем из определенных данных систем, как это было показано в приведенных примерах. Их можно также использовать для описания ситуации, с которой сталкивается шифровальщик противника, когда он пытается расшифровать криптограмму неизвестного типа. Фактически он расшифровывает секретную систему типа

$$T = p_1 A + p_2 B + \dots + p_r S + p' X, \quad \sum p_i = 1,$$

где $A, B, \dots, S$ в данном случае — известные типы шифров с их априорными вероятностями p_i , а $p' X$ соответствует возможности использования совершенно нового неизвестного шифра.

7. Чистые и смешанные шифры

Некоторые типы шифров, такие, как простая подстановка, транспозиция с данным периодом, система Виженера с данным периодом, система Виженера со смешанным алфавитом и т. д. (все с равновероятными ключами) обладают некоторой однородностью по отношению к ключу. Каков бы ни был ключ, процессы шифрования, дешифрирования адресатом и дешифрирования противником являются по существу теми же самыми. Эти системы можно противопоставить системе с шифром

$$pS + qT,$$

где S — простая подстановка, а T — транспозиция с данным периодом. В таком случае процессы шифрования и дешифрирования адресатом или противником полностью меняются в зависимости от того, используется подстановка или транспозиция.

Причина однородности таких систем лежит в групповом свойстве: заметим, что в приведенных выше примерах однородных шифров произведение $T_i T_j$ любых двух отображений из множества равно третьему отображению T_k из этого же множества. С другой стороны, $T_i S_j$ не равно какому-нибудь отображению для шифра

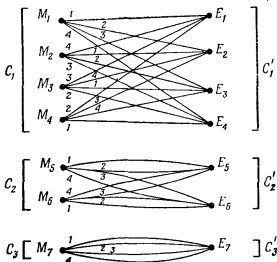
$$pS + qT,$$

который содержит только подстановки и транспозиции, но не их произведения.

Было бы можно, таким образом, определить «чистый» шифр как шифр, в котором T_i образуют группу. Однако это было бы слишком сильным ограничением, так как тогда потребовалось бы, чтобы

Остаточные
классы
сообщений

Остаточные
классы
криптограмм



Чистая система

Р и с. 4. Чистая система.

пространство E совпадало с пространством M , т. е. чтобы система была эндоморфной. Дробная транспозиция так же однородна, как и обычная транспозиция, но она не эндоморфна. Подходящим является следующее определение: шифр T является чистым, если для каждого T_i , T_j , T_k имеется такое T_s , что

$$T_i T_j^{-1} T_k = T_s,$$

и все ключи равновероятны. В противном случае шифр является смешанным. Шифры на рис. 2 являются смешанными, а на рис. 4 — чистыми, если только все ключи равновероятны.

Теорема 1. В чистом шифре операции $T_i^{-1} T_j$, отображающие пространство сообщений в себя, образуют группу, порядок которой равен m — числу различных ключей.

Так как

$$T_j^{-1} T_k T_k^{-1} T_j = I,$$

то каждый элемент имеет обратный. Ассоциативный закон верен, так как это операции, а групповое свойство следует из того, что

$$T_i^{-1} T_j T_k^{-1} T_l = T_s^{-1} T_k T_k^{-1} T_l = T_s^{-1} T_l,$$

где предполагалось, что для некоторого s $T_i^{-1} T_j = T_s^{-1} T_k$.

Операция $T_i^{-1} T_j$ означает шифрование сообщения с помощью ключа j с последующим дешифрованием с помощью ключа i , что приводит нас назад к пространству сообщений. Если система T эндоморфна, т. е. T_i отображают пространство Ω_M в само себя (что имеет место для большинства шифров, в которых и пространство сообщений и пространство криптограмм состоит из последовательностей букв), и если T_i образуют группу и равновероятны, то T — чистый шифр, так как

$$T_i T_j^{-1} T_k = T_i T_r = T_s.$$

Теорема 2. Произведение двух чистых коммутирующих шифров является чистым шифром.

Если T и R коммутируют, то $T_i R_j = R_l T_m$ для любых i, j при соответствующих l, m . Тогда

$$\begin{aligned} T_i R_j (T_k R_l)^{-1} T_m R_n &= T_i R_j R_l^{-1} T_k^{-1} T_m R_n = \\ &= R_u R_v^{-1} R_w T_r T_s^{-1} T_t = R_h T_g. \end{aligned}$$

Условие коммутирования не является, однако, необходимым для того, чтобы произведение было чистым шифром.

Система, состоящая из одного ключа, т. е. из единственной определенной операции T_1 , является чистым шифром, т. е. при единственном возможном выборе индексов имеем

$$T_1 T_1^{-1} T_1 = T_1.$$

Таким образом, разложение шифра в сумму таких простых отображений представляет собой разложение в его сумму чистых шифров.

Исследование примера, приведенного на рис. 4, вскрывает некоторые свойства чистого шифра. Сообщения распадаются на определенные подмножества, которые мы будем называть *остаточными классами*, и возможные криптограммы также распадаются на соответствующие им *остаточные классы*. От каждого сообщения в любом

классе к каждой криптограмме в соответствующем классе имеется не менее одной линии и нет линий между несоответствующими классами. Число сообщений в классе является делителем полного числа ключей. Число «параллельных» линий от сообщения M к криптограмме в соответствующем классе равно числу ключей, деленному на число сообщений в классе, содержащем это сообщение (или криптограмму).

В приложении показывается, что это верно для чистых шифров и в общем случае. Резюмируя сказанное, мы имеем

Теорема 3. *В чистой системе сообщения можно разделить на множество «остаточных» классов $C_1, \dots, C_s$, а криптограммы на соответствующее множество остаточных классов $C'_1, \dots, C'_s$. Эти классы будут иметь следующие свойства:*

1. Остаточные классы сообщений взаимно исключают друг друга и содержат все возможные сообщения. Аналогичное утверждение верно и для остаточных классов криптограмм.

2. Если зашифровать любое сообщение из класса C_i с помощью любого ключа, то получится криптограмма из класса C'_i . Дешифрирование любой криптограммы из класса C'_i с помощью любого ключа приводит к сообщению из класса C_i .

3. Число сообщений в классе C_i , скажем φ_i , равно числу криптограмм в классе C'_i и является делителем k — числа ключей.

4. Каждое сообщение из класса C_i может быть зашифровано в каждую криптограмму из класса C'_i при помощи точно k/φ_i различных ключей. То же самое верно и для дешифрирования.

Смысл понятия чистый шифр (и причина для выбора такого термина) лежит в том, что в чистом шифре все ключи являются по существу одинаковыми. Какой бы ключ ни использовался для заданного сообщения, апостериорные вероятности всех сообщений будут теми же самыми¹⁾. Чтобы показать это, заметим, что два различных ключа, примененных к одному сообщению, дадут в результате две криптограммы из одного остаточного класса, скажем C'_i . Поэтому эти две криптограммы могут быть расшифрованы с помощью k/φ_i ключей в каждое из сообщений в классе C_i и больше ни в какие возможные сообщения. Так как все ключи равновероятны,

¹⁾ Речь идет об апостериорной вероятности, которую вычисляет шифровальщик противника, не знающий ключа и сообщения, на основе полученной криптограммы в предположении, что находящаяся в его распоряжении криптограмма получена из заданного сообщения при помощи заданного ключа, т. е. условная вероятность сообщения при указании криптограммы, полученной шифровальщиком, вычисляется в предположении случайности ключа и сообщения, но затем эти условные вероятности осредняются по распределению криптограмм при заданном ключе и сообщении. — *Прим. ред.*

то апостериорные вероятности различных сообщений¹⁾ равны

$$P_E(M) = \frac{P(M) \cdot P_M(E)}{P(E)} = \frac{P(M) \cdot P_M(E)}{\sum_M P(M) P_M(E)} = \frac{P(M)}{P(C_i)},$$

где M — сообщение из класса C_i , E — криптограмма из класса C_i и сумма берется по всем M из класса C_i . Если E и M не принадлежат соответствующим остаточным классам, то $P_E(M) = 0$.

Аналогично можно показать, что набор апостериорных вероятностей различных ключей всегда одинаков, но эти вероятности ставятся в соответствие ключам лишь после того, как уже использован некоторый ключ. При изменении частного ключа это множество чисел $P_E(K)$ подвергается перестановке²⁾. Иными словами, имеем:

Теорема 4. В чистой системе апостериорные вероятности различных сообщений $P_E(M)$ не зависят от выбора ключа. Апостериорные вероятности ключей $P_E(K)$ образуют один и тот же набор величин, но подвергаются перестановке в результате различных выборов ключа.

Грубо говоря, можно считать, что любой выбор ключа в чистом шифре приводит к одинаковым трудностям при дешифрировании. Поскольку различные ключи все приводят к формированию криптограмм из одного и того же остаточного класса, то все криптограммы из одного остаточного класса эквивалентны с точки зрения сложности дешифрирования — они приводят к тем же самым апостериорным вероятностям сообщений и, если учитывать перестановки, к тем же самым вероятностям ключей.

В качестве примера чистого шифра может служить простая подстановка с равновероятными ключами. Остаточный класс, соответствующий данной криптограмме E , является множеством всех криптограмм, которые могут быть получены из E с помощью операций $T_j T_k^{-1} E$. В рассматриваемом случае операция $T_j T_k^{-1}$ сама является подстановкой и поэтому любая подстановка переводит криптограмму E в другой член того же самого остаточного класса; таким образом, если криптограмма представляет собой

$$E = XCPPGCFQ,$$

то

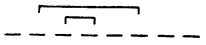
$$E_1 = RDHHGDSN,$$

$$E_2 = ABCCDBEF,$$

¹⁾ Апостериорные вероятности ключей принимаются в том же смысле, что и апостериорные вероятности сообщений в примечании на стр. 356.— *Прим. ред.*

²⁾ Здесь и в дальнейшем автор обозначает через $P_A(B)$ — условную вероятность B при условии A .— *Прим. ред.*

и т. д. принадлежат к тому же остаточному классу. В этом случае очевидно, что криптограммы по существу эквивалентны. Все существенное в простой подстановке со случайным ключом заключено в характере повторения букв, в то время как сами буквы являются несущественной маскировкой. В действительности можно бы полностью обойтись без них, указав характер повторений букв в E следующим образом:



Это обозначение описывает остаточный класс, но устраняет всю информацию относительно конкретных членов этого класса; таким образом, оно представляет как раз ту информацию, которая имеет значение для шифровальщика противника. Это связано с одним из методов подхода к раскрытию шифров типа простой подстановки — методом характерных слов.

В шифре типа Цезаря имеют значение только первые разности криптограммы по модулю 26. Две криптограммы с теми же самыми разностями (Δe_i) принадлежат к одному остаточному классу. Этот шифр можно раскрыть путем простого процесса выписывания двадцати шести сообщений из этого остаточного класса и выбора того из них, которое имеет смысл.

Шифр Виженера с периодом d со случайным ключом представляет собой другой пример чистого шифра. Здесь остаточный класс сообщений состоит из всех последовательностей с теми же первыми разностями, что и у криптограммы для букв, отстоящих на расстояние d . Для $d = 3$ остаточный класс определяется с помощью равенств

$$m_1 - m_4 = e_1 - e_4$$

$$m_2 - m_5 = e_2 - e_5$$

$$m_3 - m_6 = e_3 - e_6$$

$$m_4 - m_7 = e_4 - e_7$$

.....

где $E = e_1, e_2, \dots$ — криптограмма, а $m_1 m_2 \dots$ является любым сообщением M в соответствующем остаточном классе.

В транспозиции с периодом d со случайным ключом остаточный класс состоит из всех способов расстановок символов криптограммы, в которых никакое e_i не выдвигается из своего блока длины d , и любые два e_i с расстоянием d остаются на таком же расстоянии. Это используется для раскрытия шифра следующим образом: криптограмма записывается в виде последовательных блоков длины d

один под другим, как показано ниже (для $d = 5$)

$$\begin{array}{cccccc} e_1 & e_2 & e_3 & e_4 & e_5 & \\ e_6 & e_7 & e_8 & e_9 & e_{10} & \\ e_{11} & e_{12} & . & . & . & . \\ . & . & . & . & . & . \end{array}$$

Затем столбцы переставляются до тех пор, пока не получится осмысленный текст. После того как криптограмма разбита на столбцы, оставшейся существенной информацией является только остаточный класс криптограммы.

Теорема 5. Если шифр T — чистый, то $T_i T_j^{-1} T = T$, где T_i, T_j — любые два отображения из T . Обратно, если это выполняется для любых принадлежащих шифру T_i, T_j , то шифр T является чистым.

Первая часть этой теоремы следует, очевидно, из определения чистого шифра. Чтобы доказать вторую часть, заметим сначала, что если $T_i T_j^{-1} T = T$, то $T_i T_j^{-1} T_s$ является отображением из T . Остается показать, что все ключи равновероятны. Имеем $T = \sum_s p_s T_s$ и

$$\sum_s p_s T_i T_j^{-1} T_s = \sum_s p_s T_s.$$

Слагаемое в стоящей слева сумме с $s = j$ дает $p_j T_i$. Единственным слагаемым с T_i в правой части является $p_i T_i$. Так как все коэффициенты неотрицательны, то отсюда следует, что

$$p_j \leq p_i.$$

То же самое рассуждение остается справедливым, если i и j поменять местами. Следовательно,

$$p_i = p_j$$

и T — чистый шифр. Таким образом, условие $T_i T_j^{-1} T = T$ можно было бы использовать в качестве другого определения чистого шифра.

8. Подобные системы

Две секретные системы R и S будем называть *подобными*, если существует отображение A , имеющее обратное A^{-1} , такое, что

$$R = AS.$$

Это означает, что шифрование с помощью R даст то же, что шифрование с помощью S с последующим применением отображения A . Если использовать запись $R \approx S$ для обозначения того, что R

подобно S , то, очевидно, из $R \approx S$ следует $S \approx R$. Кроме того, из $R \approx S$ и $S \approx T$ следует, что $R \approx T$ и, наконец, $R \approx R$. Резюмируя вышеизложенное, можно сказать, что подобие систем является соотношением эквивалентности.

Криптографический смысл подобия состоит в том, что если $R \approx S$, то R и S — эквивалентны с точки зрения дешифрирования. Действительно, если шифровальщик противника перехватывает криптограмму из системы S , он может перевести ее в криптограмму из системы R простым применением к ней отображения A . Обратно, криптограмма из системы R переводится в криптограмму из системы S с помощью A^{-1} . Если R и S применяются к одному и тому же пространству сообщений или языку, то имеется взаимнооднозначное соответствие между получающимися криптограммами. Соответствующие друг другу криптограммы дают одинаковое апостериорное распределение вероятностей для всех сообщений.

Если имеется некоторый способ раскрытия системы R , то любая система S , подобная R , может быть раскрыта после приведения ее к R с помощью операции A . Этот способ часто используется на практике.

В качестве тривиального примера рассмотрим простую подстановку, в которой буквы сообщения заменяются не буквами, а произвольными символами. Она подобна обычной простой подстановке с заменой на буквы. Вторым примером могут служить шифр Цезаря и обратный шифр Цезаря. Последний иногда раскрывают, переводя его сначала в шифр Цезаря. Это можно сделать, обратив алфавит в криптограмме. Шифры Виженера, Бофора и вариант Бофора все подобны, если ключ является случайным. Шифр с «автоключом» (т. е. сообщением, используемым в качестве «ключа») с используемыми вначале ключами $K_1 K_2 \dots K_d$ подобен шифру Виженера с ключом, поочередно складываемым и вычитаемым по модулю 26. Отображение A в этом случае представляет собой «дешифровку» автоключа с помощью последовательности из d таких отображений для каждого из начальных ключей.

Часть II

ТЕОРЕТИЧЕСКАЯ СЕКРЕТНОСТЬ

9. Введение

Рассмотрим вопросы, связанные с «теоретической секретностью» систем. Насколько устойчива некоторая система, если шифровальщик противника не ограничен временем и обладает всеми необходимыми средствами для анализа криптограмм? Имеет ли криптограмма

единственное решение (даже если для нахождения этого решения может потребоваться такой объем работ, что его практически нельзя будет выполнить), а если нет, то сколько она имеет приемлемых решений? Какой объем текста, зашифрованного в данной системе, нужно перехватить для того, чтобы решение стало единственным? Существуют ли секретные системы, в которых вообще нельзя найти единственного решения независимо от того, каков объем перехваченного зашифрованного текста? Существуют ли секретные системы, в которых противник не получает никакой информации, сколько бы он ни перехватывал зашифрованного текста? В анализе этих вопросов найдут широкое применение понятия энтропии, избыточности, а также и другие понятия, введенные в работе «Математическая теория связи»¹⁾

10. Совершенная секретность

Предположим, что имеется конечное число возможных сообщений $M_1, \dots, M_n$ с априорными вероятностями $P(M_1), \dots, P(M_n)$ и что эти сообщения преобразуются в возможные криптограммы $E_1, \dots, E_m$, так что

$$E = T_i M.$$

После того как шифровальщик противника перехватил некоторую криптограмму E , он может вычислить, по крайней мере в принципе, апостериорные вероятности различных сообщений $P_E(M)$. Естественно определить *совершенную секретность* с помощью следующего условия: для всех E апостериорные вероятности равны априорным вероятностям независимо от величины этих последних. В этом случае перехват сообщения не дает шифровальщику противника никакой информации²⁾. Теперь он не может корректировать никакие свои действия в зависимости от информации, содержащейся в криптограмме, так как все вероятности, относящиеся к содержанию криптограммы, не изменяются. С другой стороны, если это условие равенства вероятностей не выполнено, то имеются такие случаи, в которых для определенного ключа и определенных выборов сообщений апостериорные вероятности противника отличаются от априорных. А это в свою очередь может повлиять на выбор

¹⁾ Перевод этой работы помещен в данном сборнике (стр. 243). — *Прим. ред.*

²⁾ Пурист мог бы возразить, что противник получил некоторую информацию, а именно он знает, что послано какое-то сообщение. На это можно ответить следующим образом. Пусть среди сообщений имеется «чистый бланк», соответствующий «отсутствию сообщения». Если не создается никакого сообщения, то чистый бланк зашифровывается и посылается в качестве криптограммы. Тогда устраняется даже эта крупинка информации.

противником своих действий и, таким образом, совершенной секретности не получится. Следовательно, приведенное определение неизбежным образом следует из нашего интуитивного представления о совершенной секретности.

Необходимое и достаточное условие для того, чтобы система была совершенно секретной, можно записать в следующем виде. По теореме Байеса

$$P_E(M) = \frac{P(M) \cdot P_M(E)}{P(E)},$$

где

- $P(M)$ — априорная вероятность сообщения M ;
- $P_M(E)$ — условная вероятность криптограммы E при условии, что выбрано сообщение M , т. е. сумма вероятностей всех тех ключей, которые переводят сообщение M в криптограмму E ;
- $P(E)$ — вероятность получения криптограммы E ;
- $P_E(M)$ — апостериорная вероятность сообщения M при условии, что перехвачена криптограмма E .

Для совершенной секретности системы величины $P_E(M)$ и $P(M)$ должны быть равны для всех E и M . Следовательно, должно быть выполнено одно из равенств: или $P(M) = 0$ [это решение должно быть отброшено, так как требуется, чтобы равенство осуществлялось при любых значениях $P(M)$], или же

$$P_M(E) = P(E)$$

для любых M и E .

Наоборот, если $P_M(E) = P(E)$, то

$$P_E(M) = P(M),$$

и система совершенно секретна. Таким образом, можно сформулировать следующее:

Теорема 6. *Необходимое и достаточное условие для совершенной секретности состоит в том, что*

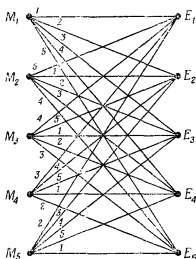
$$P_M(E) = P(E)$$

для всех M и E , т. е. $P_M(E)$ не должно зависеть от M .

Другими словами, полная вероятность всех ключей, переводящих сообщение M_i в данную криптограмму E , равна полной вероятности всех ключей, переводящих сообщение M_j в ту же самую криптограмму E для всех M_i, M_j и E .

Далее, должно существовать по крайней мере столько же криптограмм E , сколько и сообщений M , так как для фиксированного i отображение T_i дает взаимнооднозначное соответствие между всеми M и некоторыми из E . Для совершенно секретных систем для

каждого из этих E и любого M $P_M(E) = P(E) \neq 0$. Следовательно, найдется по крайней мере один ключ, отображающий данное M в любое из E . Но все ключи, отображающие фиксированное M в различные E , должны быть различными, и поэтому число различных ключей не меньше числа сообщений M . Как показывает следующий



Р и с. 5. Совершенная система.

пример, можно получить совершенную секретность, когда число сообщений точно равно числу ключей. Пусть M_i занумерованы числами от 1 до n , так же как и E_i , и пусть используются n ключей. Тогда

$$T_i M_j = E_s,$$

где $s = i + j \pmod{n}$. В этом случае оказывается справедливым равенство $P_E(M) = \frac{1}{n} = P(E)$ и система является совершенно секретной. Один пример такой системы показан на рис. 5, где

$$s = j + i - 1 \pmod{5}.$$

Совершенно секретные системы, в которых число криптограмм равно числу сообщений, а также числу ключей, характеризуются следующими двумя свойствами: 1) каждое M связывается с каждым E только одной линией; 2) все ключи равновероятны. Таким образом, матричное представление такой системы является «латинским квадратом».

В «Математической теории связи» показано, что количественно информацию удобно измерять с помощью энтропии. Если имеется некоторая совокупность возможностей с вероятностями $p_1, \dots, p_n$,

то энтропия дается выражением

$$H = - \sum p_i \log p_i.$$

Секретная система включает в себя два статистических выбора: выбор сообщения и выбор ключа. Можно измерять количество информации, создаваемой при выборе сообщения, через $H(M)$

$$H(M) = - \sum P(M) \log P(M),$$

где суммирование выполняется по всем возможным сообщениям. Аналогично, неопределенность, связанная с выбором ключа, дается выражением

$$H(K) = - \sum P(K) \log P(K).$$

В совершенно секретных системах описанного выше типа количество информации в сообщении равно самое большое $\log n$ (эта величина достигается для равновероятных сообщений). Эта информация может быть скрыта полностью лишь тогда, когда неопределенность ключа не меньше $\log n$. Это является первым примером общего принципа, который будет часто встречаться ниже: существует предел, которого нельзя превзойти при заданной неопределенности ключа — количество неопределенности, которое может быть введено в решение, не может быть больше чем неопределенность ключа.

Положение несколько усложняется, если число сообщений бесконечно. Предположим, например, что сообщения порождаются соответствующим марковским процессом в виде бесконечной последовательности букв. Ясно, что никакой конечный ключ не даст совершенной секретности. Предположим тогда, что источник ключа порождает ключ аналогичным образом, т. е. как бесконечную последовательность символов. Предположим далее, что для шифрования и дешифрирования сообщения длины L_M требуется только определенная длина ключа L_K . Пусть логарифм числа букв в алфавите сообщений будет R_M , а такой же логарифм для ключа — R_K . Тогда из рассуждений для конечного случая, очевидно, следует, что для совершенной секретности требуется, чтобы выполнялось неравенство

$$R_M L_M \leq R_K L_K.$$

Такой вид совершенной секретности реализован в системе Вернама.

Эти выводы делаются в предположении, что априорные вероятности сообщений неизвестны или произвольны. В этом случае ключ, требуемый для того, чтобы имела место совершенная секретность, зависит от полного числа возможных сообщений.

Можно было бы ожидать, что если в пространстве сообщений имеются фиксированные известные статистические связи, так что имеется определенная скорость создания сообщений R в смысле,

принятом в «Математической теории связи», то необходимый объем ключа можно было бы снизить в среднем в R/R_M раз, и это действительно верно. В самом деле, сообщение можно пропустить через преобразователь, который устраняет избыточность и уменьшает среднюю длину сообщения как раз во столько раз. Затем к результату можно применить шифр Вернама. Очевидно, что объем ключа, используемого на букву сообщения, статистически уменьшается на множитель R/R_M , и в этом случае источник ключа и источник сообщений в точности согласован — один бит ключа полностью скрывает один бит информации сообщения. С помощью методов, использованных в «Математической теории связи» легко также показать, что это лучшее, чего можно достигнуть.

Совершенно секретные системы могут применяться и на практике их можно использовать или в том случае, когда полной секретности придается чрезвычайно большое значение, например для кодирования документов высших военных инстанций управления, или же в случаях, где число возможных сообщений мало. Так, беря крайний пример, когда имеются в виду только два сообщения — «да» или «нет», — можно, конечно, использовать совершенно секретную систему со следующей таблицей отображений:

$M \ K$	$A \ B$
да	0 1
нет	1 0

Недостатком совершенно секретных систем для случая корреспонденции большого объема является, конечно, то, что требуется посылать эквивалентный объем ключа. В следующих разделах будет рассмотрен вопрос о том, чего можно достигнуть при помощи меньших объемов ключа, в частности с помощью конечного ключа.

11. Ненадежность

Предположим теперь, что для английского текста используется шифр простой подстановки и что перехвачено определенное число, скажем N , букв зашифрованного текста. Если N достаточно велико, скажем более 50, то почти всегда существует единственное решение шифра, т. е. единственная последовательность, имеющая смысл на английском языке, в которую переводится перехваченный материал с помощью простой подстановки. Для меньших N шансы на неединственность решения увеличиваются; для $N = 15$, вообще говоря, будет существовать некоторое число подходящих отрывков осмысленного английского текста, в то время как для $N = 8$ ока-

жется подходящей значительная часть (порядка $\frac{1}{8}$) всех возможных значащих английских последовательностей такой длины, так как из восьми букв редко повторится больше чем одна. При $N = 1$ очевидно, возможна любая буква и апостериорная вероятность любой буквы будет равна ее априорной вероятности. Для одной буквы система является совершенно секретной.

Это происходит, вообще говоря, со всеми разрешимыми шифрами. Прежде чем перехвачена криптограмма, можно представить себе априорные вероятности, связанные с различными возможными сообщениями, а также с различными ключами. После того как материал перехвачен, шифровальщик противника вычисляет их апостериорные вероятности. При увеличении числа N вероятности некоторых сообщений возрастают, но для большинства сообщений они убывают до тех пор, пока не останется только одно сообщение, имеющее вероятность, близкую к единице, в то время как полная вероятность всех других близка к нулю.

Для самых простых систем эти вычисления можно эффективно выполнить. Таблица I дает апостериорные вероятности для шифра Цезаря, примененного к английскому тексту, причем ключ выбирался случайно из 26 возможных ключей. Для того чтобы можно было использовать обычные таблицы частот букв, диграмм и триграмм, текст был начат в случайном месте (на страницу открытой наугад книги был случайно опущен карандаш). Сообщение, выбранное таким способом, начинается с «creases to» (карандаш опущен на третью букву слова increases). Если известно, что сообщение начинается не с середины, а с начала некоторого предложения, то нужно пользоваться иной таблицей, соответствующей частотам букв, диграмм и триграмм, стоящих в начале предложения.

Шифр Цезаря со случайным ключом является чистым, и выбор частного ключа не влияет на апостериорные вероятности. Чтобы определить эти вероятности, надо просто выписать возможные расшифровки с помощью всех ключей и вычислить их априорные вероятности. Апостериорные вероятности получатся из этих последних в результате деления их на их сумму. Эти возможные расшифровки, образующие остаточный класс этого сообщения, найдены с помощью стандартного процесса последовательного «пробега алфавита», на табл. I они даны слева. Для одной перехваченной буквы апостериорные вероятности равны априорным вероятностям для всех букв¹⁾ (они приведены на таблице под рубрикой $N = 1$).

¹⁾ Вероятности в приводимой таблице были взяты из таблиц частот, данных в книге Pratt F., *Secret and Urgent*, Blue Ribbon Books, New York, 1939. Хотя эти таблицы и не являются полными, но для настоящих целей их достаточно.

Таблица 1

Апостериорные вероятности для криптограммы типа Цезаря

Расшифровки	$N = 1$	$N = 2$	$N = 3$	$N = 4$	$N = 5$
CREAS	0,028	0,0377	0,1111	0,3673	1
DSFBT	0,038	0,0314			
ETGCU	0,131	0,0881			
FUNDV	0,029	0,0189			
GVIEW	0,020				
HWJFX	0,053	0,0063			
IXKGY	0,063	0,0126			
JYLHZ	0,001				
KZMIA	0,004				
LANJB	0,034	0,1321	0,2500		
MBOKC	0,025		0,0222		
NCPLD	0,071	0,1195			
ODQME	0,080	0,0377			
PERNF	0,020	0,0818	0,4389	0,6327	
QFSOG	0,001				
RGTPH	0,068	0,0126			
SHUQI	0,061	0,0881	0,0056		
TIVRJ	0,105	0,2830	0,1667		
UJWSK	0,025				
VKXTL	0,009				
WLYUM	0,015		0,0056		
XMZVN	0,002				
YNAWO	0,020				
ZOBXP	0,001				
APCYQ	0,082	0,0503			
BQDZR	0,014				
H (десятичных единиц)	1,2425	0,9686	0,6034	0,285	0

Для двух перехваченных букв эти вероятности равны априорным вероятностям диграмм, пронормированным на их сумму (они приведены в столбце $N = 2$). Триграммные частоты получены аналогично и приведены в столбце $N = 3$. Для четырех- и пятибуквенных последовательностей вероятности находились из триграммных частот с помощью умножения, так как с некоторым приближением

$$p(ijkl) = p(ijk) p_{jk}(l).$$

Заметим, что для трех букв число возможных сообщений снижается до четырех сообщений достаточно высокой вероятности, причем вероятности всех других сообщений малы по сравнению с вероятностями этих четырех сообщений. Для четырех букв имеются два возможных сообщения и для пяти — только одно, а именно правильная дешифровка.

В принципе это может быть проведено для любой системы, однако в том случае, когда объем ключа не очень мал, число возможных сообщений настолько велико, что вычисления становятся практически невыполнимыми.

Получаемое таким образом множество апостериорных вероятностей описывает, как постепенно по мере получения зашифрованного материала становятся более точными сведения о шифровальщике противника относительно сообщения и ключа.

Это описание, однако, является слишком исчерпывающим и слишком сложным для наших целей. Хотелось бы иметь упрощенное описание такого приближения к единственности возможного решения.

Аналогичная ситуация возникает в теории связи, когда передаваемый сигнал искажается шумом. Здесь необходимо ввести подходящую меру неопределенности того, что действительно было передано, при условии, что известен только искаженный шумом вариант — принятый сигнал.

В «Математической теории связи» показано, что естественной математической мерой этой неопределенности является условная энтропия передаваемого сигнала при условии, что принятый сигнал известен. Эта условная энтропия для удобства будет называться ненадежностью.

С криптографической точки зрения секретная система почти тождественна системе связи при наличии шума. На сообщение (передаваемый сигнал) действует некоторый статистический элемент (секретная система с ее статистически выбранным ключом). В результате получается криптограмма (аналог искаженного сигнала), подлежащая дешифрированию. Основное различие заключается в следующем: во-первых, в том, что преобразование при помощи шифра имеет обычно более сложную природу, чем возникающее за счет шума в канале; и, во-вторых, ключ в секретной системе обычно выбирается из конечного множества, в то время как шум в канале чаще является непрерывным, выбранным по существу из бесконечного множества.

Учитывая эти соображения, естественно использовать ненадежность в качестве теоретической меры секретности. Следует отметить, что имеются две основные ненадежности: ненадежность ключа и ненадежность сообщения. Они будут обозначаться через $H_E(K)$ и $H_E(M)$ соответственно. Их величины определяются

соотношениями

$$H_E(K) = - \sum_{E, K} P(E, K) \log P_E(K),$$

$$H_E(M) = - \sum_{E, M} P(E, M) \log P_E(M),$$

где E , M и K — криптограмма, сообщение и ключ;

$P(E, K)$ — вероятность ключа K и криптограммы E ;

$P_E(K)$ — апостериорная вероятность ключа K , если перехвачена криптограмма E ;

$P(E, M)$ и $P_E(M)$ — аналогичные вероятности, но не для ключа, а для сообщения.

Суммирование в $H_E(K)$ проводится по всем возможным криптограммам определенной длины (скажем, N) и по всем возможным ключам. Для $H_E(M)$ суммирование проводится по всем сообщениям и криптограммам длины N . Таким образом, $H_E(K)$ и $H_E(M)$ являются функциями от N — числа перехваченных букв. Это будет иногда указываться в обозначении так: $H_E(K, N)$, $H_E(M, N)$. Заметим, что эти ненадежности являются «полными», т. е. не делятся на N с тем, чтобы получить скорость ненадежности, которая рассматривалась в работе «Математическая теория связи».

Те же самые рассуждения, которые были использованы в «Математической теории связи» для обоснования введения ненадежности в качестве меры неопределенности в теории связи, применимы и здесь. Так, из того, что ненадежность равна нулю, следует, что одно сообщение (или ключ) имеет единичную вероятность, а все другие — нулевую. Этот случай соответствует полной осведомленности шифровальщика. Постепенное убывание ненадежности с ростом N соответствует увеличению сведений об исходном ключе или сообщении. Кривые ненадежности сообщения и ключа, нанесенные на график как функции от N , мы будем называть характеристиками ненадежности рассматриваемой секретной системы.

Величины $H_E(K, N)$ и $H_E(M, N)$ для криптограммы шифра Цезаря, рассмотренной выше, сосчитаны и приведены в нижней строке табл. I. Числа $H_E(K, N)$ и $H_E(M, N)$ в этом случае равны и даны в десятичных единицах (т. е. при вычислениях в качестве основания логарифма бралось 10). Следует отметить, что ненадежность здесь сосчитана для частной криптограммы, так как суммирование ведется только по M (или K), но не по E . В общем случае суммирование должно было бы проводиться по всем перехваченным криптограммам длины N , в результате чего получилась бы средняя неопределенность. Вычислительные трудности не позволяют сделать это практически.

12. Свойства ненадежности

Можно показать, что ненадежность обладает некоторыми интересными свойствами, большинство из которых соответствует нашему интуитивному представлению о поведении величины такого рода. Покажем сначала, что ненадежность ключа или фиксированной части сообщения уменьшается при увеличении количества перехваченного зашифрованного текста.

Теорема 7. *Ненадежность ключа $H_E(K, N)$ — невозрастающая функция N ¹⁾. Ненадежность первых A букв сообщения является невозрастающей функцией N . Если перехвачено N букв, то ненадежность первых N букв сообщения меньше или равна ненадежности ключа. Это можно записать следующим образом:*

$$H_E(K, S) \leq H_E(K, N), \quad S \geq N,$$

$$H_E(M, S) \leq H_E(M, N),$$

$$H_E(M, N) \leq H_E(K, N).$$

Введенное во втором утверждении ограничение A буквами означает, что ненадежность вычисляется по отношению к первым буквам сообщения, а не ко всему объему перехваченного сообщения. Если отказаться от этого ограничения, то можно получить возрастание ненадежности сообщения (и обычно это имеет место) с увеличением времени просто из-за того, что большее количество букв допускает и большее разнообразие возможных сообщений. Выводы этой теоремы соответствуют тому, на что можно было бы надеяться при разумной мере секретности, так как едва ли можно оказаться в худшем положении при увеличении объема перехваченного текста. Тот факт, что эти выводы могут быть доказаны, дает лучшее подтверждение полезности принятой нами количественной меры ненадежности.

Справедливость утверждений этой теоремы вытекает из некоторых свойств условной энтропии, доказанных в работе «Математическая теория связи». Так, для доказательства первого или второго утверждения теоремы воспользуемся тем, что для любых случайных событий A и B

$$H(B) \geq H_A(B).$$

Если отождествить B с ключом (при условии, что известны первые S букв криптограммы), а A с остающимися $N - S$ буквами, то мы получим первое утверждение. Аналогично, если отождествить B с сообщением, то получится второе утверждение. Последнее утвер-

¹⁾ Здесь предполагается, что ключ фиксирован и не зависит от длины криптограммы N и длины сообщения A . — Прим. ред.

ждение следует из неравенства

$$H_E(M) \leq H_E(K, M) = H_E(K) + H_{E, K}(M)$$

и из того, что $H_{E, K}(M) = 0$, так как K и E полностью определяют M .

Так как сообщение и ключ выбираются независимо, то

$$H(M, K) = H(M) + H(K).$$

Кроме того,

$$H(M, K) = H(E, K) = H(E) + H_E(K),$$

что вытекает из того факта, что знание M и K или E и K эквивалентно знанию всех трех величин M , K и E . Преобразуя эти две формулы, мы получаем формулу для ненадежности ключа:

$$H_E(K) = H(M) + H(K) - H(E).$$

В частности, если $H(M) = H(E)$, то ненадежность ключа $H_E(K)$ равна априорной неопределенности ключа $H(K)$. Это имеет место в совершенно секретных системах, описанных выше.

Формула для ненадежности сообщения может быть получена аналогичным способом. Мы имеем:

$$H(M, E) = H(E) + H_E(M) = H(M) + H_M(E);$$

$$H_E(M) = H(M) + H_M(E) - H(E).$$

Если имеется произведение секретных систем $S = TR$, то следует ожидать, что повторный процесс шифрования не уменьшит ненадежности сообщения. То, что это действительно так, можно показать следующим образом. Пусть M , E_1 , E_2 — сообщение и первая и вторая криптограммы соответственно. Тогда

$$P_{E_1 E_2}(M) = P_{E_1}(M).$$

Следовательно,

$$H_{E_1, E_2}(M) = H_{E_1}(M);$$

так как для любых случайных величин x , y , z справедливо $H_{xy}(z) \leq H_y(z)$, то получаем желаемый результат: $H_{E_2}(M) \geq H_{E_1}(M)$.

Теорема 8. *Ненадежность сообщения для произведения секретных систем $S = TR$ не меньше ненадежности для одной системы R .*

Предположим, что имеется система T , которая может быть записана как взвешенная сумма нескольких систем R , S , ..., U

$$T = p_1 R + p_2 S + \dots + p_m U, \quad \sum p_i = 1,$$

и системы R , S , ..., U имеют ненадежности $H_1, H_2, \dots, H_m$.

Теорема 9. *Ненадежность для взвешенной суммы систем ограничена неравенствами*

$$\sum p_i H_i \leq H \leq \sum p_i H_i - \sum p_i \log p_i.$$

Эти границы нельзя улучшить. Здесь H_i могут означать ненадежность ключа или сообщения.

Верхняя граница достигается, например, в строго идеальных системах (которые будут описаны ниже), где разложение производится на простые преобразования системы. Нижняя граница достигается, если все системы $R, S, \dots, U$ приводят к полностью различным пространствам криптограмм. Эта теорема также доказывается с помощью общих неравенств, которым подчиняется ненадежность:

$$H_A(B) \leq H(B) \leq H(A) + H_A(B),$$

где A может обозначать данную используемую систему, а B — ключ или сообщение.

Имеется аналогичная теорема для взвешенных сумм языков. Для ее доказательства обозначим данный язык буквой A .

Теорема 10. *Предположим, что система может быть применена к языкам $L_1, L_2, \dots, L_m$ и при этом получают ненадежности $H_1, H_2, \dots, H_m$. Если система применяется к взвешенной сумме $\sum p_i L_i$, то ненадежность H ограничена неравенствами*

$$\sum p_i H_i \leq H \leq \sum p_i H_i - \sum p_i \log p_i.$$

Эти границы нельзя улучшить. Рассматриваемая ненадежность может относиться как к ключу, так и к сообщению.

Полная избыточность D_N для N букв сообщения определяется с помощью соотношения

$$D_N = \log G - H(M),$$

где G — полное число сообщений длины N , а $H(M)$ — неопределенность выбора одного из них. В секретной системе, где полное число возможных криптограмм равно числу возможных сообщений длины N , имеет место неравенство $H(E) \leq \log G$. Следовательно,

$$\begin{aligned} H_E(K) &= H(K) + H(M) - H(E) \geq \\ &\geq H(K) - [\log G - H(M)]. \end{aligned}$$

Поэтому

$$H(K) - H_E(K) \leq D_N.$$

Из этого видно, что, например, в замкнутой системе уменьшение ненадежности ключа после перехвата N букв не превзойдет избыточности N букв языка. В таких системах (к ним относится боль-

шинство шифров) только наличие избыточности в исходном сообщении и дает возможность нахождения решения.

Предположим теперь, что имеется чистая секретная система. Обозначим различные остаточные классы сообщений через $C_1, C_2, \dots, C_r$ и соответствующие остаточные классы криптограмм через $C'_1, C'_2, \dots, C'_r$. Вероятности всех E из C'_i одинаковы

$$P(E) = \frac{P(C_i)}{\varphi_i}, \quad E - \text{элемент } C_i,$$

где φ_i — число различных сообщений в C_i . Таким образом, имеем

$$\begin{aligned} H(E) &= - \sum_i \varphi_i \frac{P(C_i)}{\varphi_i} \log \frac{P(C_i)}{\varphi_i} = \\ &= - \sum_i P(C_i) \log \frac{P(C_i)}{\varphi_i}. \end{aligned}$$

Подставив это значение $H(E)$ в выражение, полученное выше для $H_E(K)$, получим следующую теорему.

Теорема 11. Для чистого шифра

$$H_E(K) = H(K) + H(M) + \sum_i P(C_i) \log \frac{P(C_i)}{\varphi_i}.$$

Это выражение может быть использовано для вычисления $H_E(K)$ в некоторых случаях, представляющих интерес.

13. Ненадежность простой подстановки для языка с двухбуквенным алфавитом

Подсчитаем теперь ненадежность ключа и сообщения для простой подстановки¹⁾, примененной к языку с двухбуквенным алфавитом, причем вероятности для 0 и 1 равны p и q , а последовательные буквы выбираются независимо.

В этом случае

$$H_E(M) = H_E(K) = - \sum P(E) P_E(K) \log P_E(K).$$

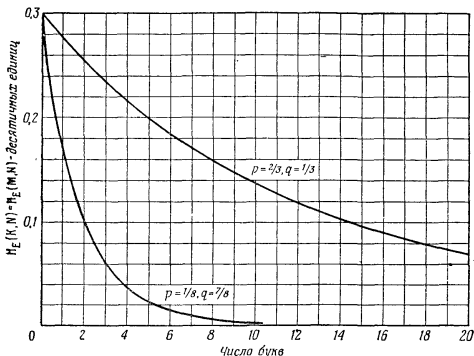
Вероятность того, что E содержит точно s нулей в фиксированных местах, равна

$$\frac{1}{2} (p^s q^{N-s} + q^s p^{N-s}),$$

и апостериорные вероятности тождественной и обратной подстановок (здесь есть только эти две подстановки) равны соответственно

$$P_E(0) = \frac{p^s q^{N-s}}{(p^s q^{N-s} + q^s p^{N-s})}, \quad P_E(1) = \frac{p^{N-s} q^s}{(p^s q^{N-s} + q^s p^{N-s})}.$$

¹⁾ В которой оба ключа (обе подстановки) равновероятны. — Прим. ред.



Р и с. 6. Ненадежность для простой подстановки в двухбуквенном языке.

Имеется C_N^s слагаемых для каждого s и, следовательно,

$$H_E(K, N) = - \sum_s C_N^s p^s q^{N-s} \log \frac{p^s q^{N-s}}{p^s q^{N-s} + q^s p^{N-s}}.$$

Для $p = \frac{1}{3}$; $q = \frac{2}{3}$ и $p = \frac{1}{8}$; $q = \frac{7}{8}$ величины $H_E(K, N)$ приведены на рис. 6.

14. Характеристика ненадежности для «случайного» шифра

В предыдущем разделе была вычислена характеристика ненадежности для простой подстановки, примененной к языку с двухбуквенным алфавитом. Но даже для этого почти самого простого шифра и языка полученные формулы уже настолько сложны, что почти бесполезны. Что же делать в практически интересных случаях, скажем в тех случаях, когда дробная транспозиция применена к английскому тексту с его крайне сложной статистической структурой. Эта сложность сама выдвигает метод подхода. Достаточно сложные проблемы часто могут быть разрешены статистически. Чтобы облегчить дело, введем понятие «случайного» шифра.

Сделаем следующие допущения.

1. Число возможных сообщений длины N равно $T = 2^{R_0 N}$; таким образом $R_0 = \log_2 G$, где G — число букв алфавита. Предполагается, что число возможных криптограмм длины N также равно T .

2. Все возможные сообщения длины N можно разделить на две группы: группу с высокими и достаточно равномерными априорными вероятностями и группу с пренебрежимо малой полной вероятностью. Высоковероятная группа будет содержать $S = 2^{RN}$ сообщений, где $R = H(M)/N$, т. е. R — энтропия источника сообщений на одну букву¹).

3. Операцию дешифрирования можно представлять графически в виде ряда линий (как на рис. 2 и рис. 4), идущих от каждого E к различным M . Предположим, что имеется k равновероятных ключей и что от каждого E будет отходить k линий. Предположим, что для случайного шифра линии от каждого E отходят к случайному набору возможных сообщений. Тогда случайный шифр будет представлять собой фактически целый ансамбль шифров и его ненадежность будет равна средней ненадежности этого ансамбля²).

Ненадежность ключа определяется с помощью равенства

$$H_E(K) = \sum P(E) P_E(K) \log P_E(K).$$

Вероятность того, что от частного E к высоковероятной группе сообщений отходит ровно m линий, равна

$$C_k^m (S/T)^m (1 - S/T)^{k-m}.$$

Если перехвачена криптограмма, которой соответствует m таких линий, то ненадежность равна $\log m$. Вероятность такой криптограммы равна mT/Sk , так как она может быть создана из высоковероятных сообщений, каждое из которых имеет вероятность T/S с помощью одного из m ключей. Отсюда ненадежность равна

$$H_E(K) = \frac{T}{Sk} \sum_{m=1}^k C_k^m (S/T)^m \left(1 - \frac{S}{T}\right)^{k-m} m \log m.$$

Требуется найти простое приближенное выражение для $H_E(K)$, когда k велико. Если для величины m ее среднее значение $\bar{m} = Sk/T$

¹) Подробнее исследование этого допущения можно найти, например, в работах А. Я. Хинчина и А. Файнштейна, указанных в библиографии в конце книги. — *Прим. ред.*

²) Это третье допущение, а также дальнейшие рассуждения этого раздела, хотя и достаточно наглядны, все же, конечно, далеки от математической строгости. Вероятно, они могут быть уточнены так же, как это уже сделано, с аналогичными построениями Шеннона в других его статьях. При этом свое место в теории могли бы занять и оговорки, которым посвящен следующий параграф. Однако эта работа, по-видимому, еще никем не проведена. — *Прим. ред.*

много больше 1, то изменение $\log m$ в области тех m , для которых биномиальные слагаемые велики, будет малым и мы можем заменить $\log m$ на $\log \bar{m}$. Этот множитель можно вынести за знак суммы, которая даст значение $\bar{m}$. Таким образом, при этом условии

$$H_E(K) \approx \log \frac{Sk}{T} = \log S - \log T + \log k,$$

$$H_E(K) \approx H(K) - DN,$$

где D — избыточность на букву первоначального языка ($D = D_N/N$).

Если $\bar{m}$ мало по сравнению с k , то биномиальное распределение может быть приближенно пуассоновским:

$$C_k^m p^m q^{k-m} \approx \frac{e^{-\lambda} \lambda^m}{m!},$$

где $\lambda = Sk/T$. Отсюда

$$H_E(K) \approx \frac{1}{\lambda} e^{-\lambda} \sum_{m=2}^{\infty} \frac{\lambda^m}{m!} m \log m.$$

Заменив m на $m+1$, получим

$$H_E(K) \approx e^{-\lambda} \sum_{m=1}^{\infty} \frac{\lambda^m}{m!} \log(m+1).$$

Это выражение можно использовать, когда λ близко к 1. Если же $\lambda \ll 1$, то существенным является лишь первый член ряда. Отбрасывая другие, получим

$$H_E(K) \approx e^{-\lambda} \lambda \log 2 \approx \lambda \log 2 \approx 2^{-ND} k \log 2.$$

Итак, подводя итог вышесказанному, получаем, что $H_E(K)$, рассматриваемая как функция от N — числа перехваченных букв — при $N = 0$ равна $H(K)$. Далее она убывает линейно с наклоном $-D$ до окрестности точки $N = H(K)/D$. После небольшой переходной области $H_E(K)$ начинает убывать экспоненциально со временем «полураспада» $1/D$, если D измеряется в битах на букву. Поведение $H_E(K)$ показано на рис. 7 вместе с аппроксимирующими кривыми.

С помощью аналогичных рассуждений можно подсчитать ненадежность сообщения. Она равна

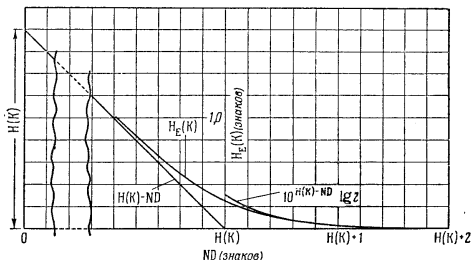
$$H_E(M) = R_0 N \quad \text{для} \quad R_0 N \ll H_E(K),$$

$$H_E(M) = H_E(K) \quad \text{для} \quad R_0 N \gg H_E(K),$$

$$H_E(M) = H_E(K) - \varphi(N) \quad \text{для} \quad R_0 N \sim H_E(K),$$

где $\varphi(N)$ — функция, показанная на рис. 7, причем шкала N сжата в D/R_0 раз. Таким образом, $H_E(M)$ возрастает линейно с наклоном R_0 до тех пор, пока она не пересечет линию $H_E(K)$. После закругленного перехода она идет ниже кривой $H_E(K)$.

На рис. 7 видно, что кривые ненадежности стремятся к нулю довольно резко. Поэтому можно с достаточной определенностью



Р и с. 7. Ненадежность для случайного шифра.

говорить о точке, в которой решение становится единственным. Найденное при этом число букв мы будем называть расстоянием единственности. Для случайного шифра оно приблизительно равно $H(K)/D$.

15. Применение к стандартным шифрам

Большинство стандартных шифров требует выполнения достаточно сложных операций шифрования и дешифрирования. Кроме того, естественные языки обладают крайне сложной статистической структурой. Поэтому можно предположить, что для случая стандартного шифра применимы формулы, выведенные для случайного шифра. Однако в некоторых случаях необходимо вносить определенные поправки. Наиболее важно отметить следующее.

1. Мы предполагали, что для случайных шифров возможные расшифровки некоторой криптограммы являются случайной выборкой из возможных сообщений. В то время как в обычных системах это, строго говоря, неверно, некоторое приближение к такой ситуации достигается при возрастании сложности операций шифрования и при усложнении статистической структуры языка. Ясно,

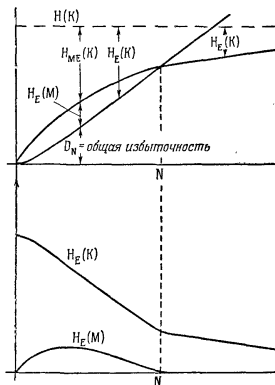
что для шифра транспозиции частоты букв сохраняются при операциях шифрования. Это означает, что возможные расшифровки выбираются не из всего пространства сообщений, а из более ограниченной группы, и, следовательно, наша формула должна быть изменена. Вместо R_0 нужно использовать R_1 — энтропийную скорость для языка с независимыми буквами, но с регулярными частотами букв. В некоторых других случаях можно заметить определенную тенденцию возвращения расшифровок к высоковероятным сообщениям. Если не заметно явно выраженной тенденции такого рода и система достаточно сложна, то разумно воспользоваться результатами анализа случайного шифра.

2. Во многих случаях полный ключ не используется при зашифровке коротких сообщений. Например, в простой подстановке все буквы алфавита будут содержаться лишь в достаточно длинных сообщениях, и таким образом, только для таких сообщений нужен полный ключ. В таком случае, очевидно, наше предположение о случайности не верно для малых N , так как все ключи, отличающиеся только в тех буквах, которые еще не появились в криптограмме, приводят к тому же самому сообщению, следовательно, эти ключи не являются случайно распределенными. Эту ошибку легко исправить, применяя «характеристику появления ключа». Для частного значения N используется эффективный объем ключа, который можно ожидать для данной длины криптограммы. Для большинства шифров этот объем легко оценить.

3. В связи с тем что сообщение начинается с определенного знака, имеют место так называемые концевые эффекты, приводящие к отклонению от случайных характеристик. Если в английском тексте взять случайную начальную точку, то первая буква (если предыдущие буквы неизвестны) может быть любой буквой с обычным набором вероятностей. Следующая буква характеризуется полнее, так как теперь имеются частоты диграмм. Это снижение в неопределенности выбора продолжается в течение некоторого времени. Действие его на кривую ненадежности выражается в том, что прямолинейная часть ее смещается и приближается к кривой, зависящей от того, в какой мере статистическая структура языка определяется соседними буквами. В качестве первого приближения кривую можно уточнить, передвигая линию до точки половины избыточности, т. е. до числа букв, для которого избыточность языка равна половине ее начального значения.

Если учитывать вышеперечисленные три особенности, то можно дать разумные оценки ненадежности и точки единственности. Вычисления можно произвести графически, как показано на рис. 8. Берется характеристика появления ключа и кривая полной избыточности D_N (которая обычно достаточно хорошо изображается прямой ND_∞). Разность между ними всюду вплоть до окрест-

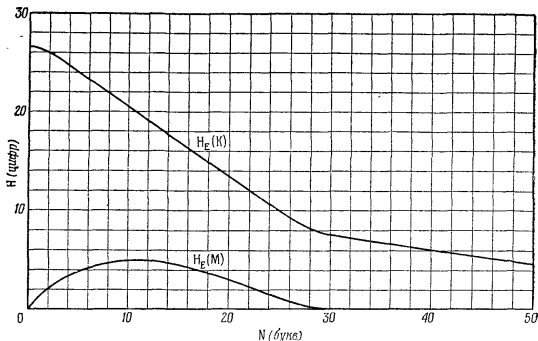
ности точки их пересечения дает $H_E(M)$. Для шифра простой подстановки, примененного к английскому тексту, это вычисление дало кривые, показанные на рис. 9. Характеристика появления ключа в этом случае оценивалась при помощи подсчета числа различных букв, появляющихся в нормативном английском отрывке из M букв. Кривые на рис. 9 очень хорошо согласуются с экспериментальными данными, в той мере, в какой они могут быть получены для простой подстановки, если учесть при этом, что были сделаны



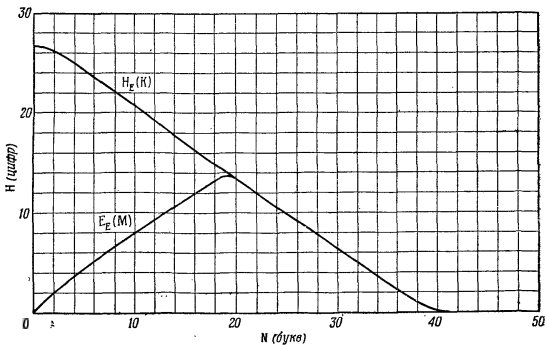
Р и с. 8. Графическое вычисление ненадежности.

различные идеализации и приближения. Например, можно показать, что точка единственности, для которой получается теоретически значение, равное примерно 27 буквам, в экспериментах заключена в пределах от 20 до 30 букв. С помощью 30 букв почти всегда получается единственное решение криптограммы такого типа, а с помощью 20 букв можно обычно легко найти несколько решений.

Для транспозиции с периодом d (при случайном ключе) $H(K) = \log d! \approx d \log(d/e)$ (приближение получено с помощью формулы Стирлинга). Если в качестве подходящей избыточности взять значение 0,6 десятичных единиц на букву (с учетом сохранения



Р и с. 9. Ненадежность для простой подстановки в английском языке.



Р и с. 10. Ненадежность для шифра Виженера для английского языка.

частот букв), то для расстояния единственности получим приблизительно величину $1,7 d \log d/e$. Эта величина также хорошо подтверждается экспериментом. Заметим, что в этом случае $H_E(M)$ определено только для целых кратных d .

Для шифра Виженера теоретически рассчитанная точка единственности лежит в окрестности $2d$ букв, и это также близко к наблюдаемым значениям. Ненадежность шифра Виженера с тем же объемом ключа, что и у простой подстановки, будет примерно такой, как показано на рис. 10. Шифры Виженера, Плайфер и дробный шифр более точно подчиняются теоретическим формулам для случайных шифров, чем простая подстановка и транспозиция. Причина этого лежит в том, что первые являются более сложными и приводят к более перемешанным характеристикам тех сообщений, к которым они применены.

Шифр Виженера с перемешанным алфавитом (каждый из d алфавитов перемещивается независимо, и алфавиты используются последовательно) имеет объем ключа, равный

$$H(K) = d \log 26! = 26,3d,$$

и точка единственности должна лежать около $53d$ букв.

Эти выводы могут быть использованы также для грубой экспериментальной проверки теории для шифров типа Цезаря. Для конкретной криптограммы, проанализированной в табл. 1 разд. 11, функция $H_E(K, N)$ вычислена и приведена ниже вместе с аналогичными величинами для случайного шифра

N	0	1	2	3	4	5
H (наблюденная)	1,41	1,24	0,97	0,60	0,28	0
H (вычисленная)	1,41	1,25	0,98	0,54	0,15	0,03

Согласие, как видно, достаточно хорошее, особенно если учесть, что наблюдаемая H должна быть в действительности средней для многих различных криптограмм и что для больших N величина D оценивается лишь весьма грубо.

Таким образом, оказывается, что методы анализа случайного шифра могут быть использованы для оценки характеристик ненадежности и расстояния единственности обычных типов шифров.

16. Правильность решения криптограммы

Формулы для расчета ненадежности относятся и к вопросам, которые иногда возникают в криптографических работах, при рассмотрении правильности предлагаемого решения криптограммы.

В истории криптографии зарегистрировано много криптограмм или возможных криптограмм, для которых искусные специалисты находили «решение». Однако это требовало выполнения таких сложных преобразований или выполнялось на таком скудном материале, что возникал вопрос не «приписал» ли шифровальщик смысл анализируемой криптограмме. Например, шифры Бэкона — Шекспира и рукопись Роджера Бэкона¹⁾.

В общем случае можно сказать, что если предлагаемая система и ключ решают криптограмму для количества материала, которое значительно превосходит расстояние единственности, то решение заслуживает доверия. Если же количество материала равно (или меньше) расстояния единственности, то правильность решения весьма сомнительна.

Действие избыточности при постепенном получении единственного решения полезно представить себе следующим образом. Избыточность представляет собой по существу ряд условий, наложенных на буквы сообщения, которые обеспечивают его надлежащую статистику. Эти ограничительные условия создают соответствующие ограничительные условия в криптограмме. Ключ создает некоторую свободу в решении криптограммы, но по мере того, как перехватываются очередные буквы, ограничительные условия исключают свободу, даваемую ключом. В конце концов остается только одно сообщение и один ключ, которые удовлетворяют всем условиям, и находится единственное решение. В случайном шифре ограничительные условия в некотором смысле «ортогональны» «структуре ключа» и в полной мере способствуют возможно скорейшему исключению всех сообщений и ключей, отличных от искомым. Это наблюдается в обычных случаях. Однако с помощью соответствующих систем можно «выровнять» избыточность языка и «структуру ключа», так что ограничительные условия будут удовлетворяться автоматически и $H_E(K)$ не будет стремиться к нулю. Эти «идеальные» системы, рассматриваемые в следующем разделе, таковы, что все отображения T_i приводят к одинаковым вероятностям в пространстве E .

17. Идеальные секретные системы

Как уже было показано, в совершенно секретных системах для сообщений неограниченной длины требуется ключ бесконечного объема. Если использовать ключ конечного объема, то ненадежности ключа и сообщения, вообще говоря, будут стремиться к нулю, хотя это и не обязательно. На самом деле можно удерживать значение $H_E(K)$, равным ее начальному значению $H(K)$. Тогда, независимо

¹⁾ См. работу Ф. Пратта, цитированную на стр. 366.

от того, сколько зашифрованного материала перехвачено, единственного решения не будет, а будет много решений со сравнимыми по величине вероятностями. Определим «идеальную» систему как такую, в которой величины $H_E(K)$ и $H_E(M)$ не стремятся к нулю при $N \rightarrow \infty$. «Строго идеальная» система — это такая, в которой величина $H_E(K)$ остается равной $H(K)$.

Примером последней может служить простая подстановка, примененная к искусственному языку, в котором все буквы равновероятны и последовательные буквы выбираются независимо. Легко видеть, что здесь $H_E(K) = H(K)$ и $H_E(M)$ растет линейно по прямой с наклоном $\log G$ (где G — число букв в алфавите) до тех пор, пока она не пересечет линию $H(K)$, после чего она остается равной этой константе.

Для естественных языков можно, вообще говоря, приблизиться к идеальной характеристике, т. е. отодвинуть точку единственности на сколько угодно большое расстояние. Однако если попытаться это сделать, то сложность требующейся системы будет обычно быстро возрастать. Не всегда возможно достичь идеальной характеристики с помощью какой-либо системы ограниченной сложности.

Для того чтобы приблизиться к идеальной ненадежности, можно преобразовать сообщение с помощью устройства, которое устраняет всю избыточность. После этого достаточно применить любой шифр — подстановку, транспозицию, шифр Виженера и т. д. Чем тщательнее сконструировано преобразующее устройство и чем ближе его выход к желаемой форме, тем точнее секретная система будет приближаться к идеальной.

Теорема 12. *Необходимое и достаточное условие строгой идеальности системы T заключается в том, что для любых двух ключей отображение $T_1^{-1}T_2$ должно являться сохраняющим меру отображением пространства сообщений в само себя.*

Это верно, так как апостериорная вероятность каждого ключа равна его априорной вероятности тогда и только тогда, когда выполнено это условие.

18. Примеры идеальных секретных систем

Предположим, что наш язык состоит из последовательности букв, выбираемых независимо и с равными вероятностями. Тогда избыточность равна нулю и из выводов разд. 12 следует, что $H_E(K) = H(K)$. Получаем следующий результат.

Теорема 13. *Если все буквы равновероятны и выбираются независимо, то любая замкнутая система будет строго идеальной.*

Ненадежность сообщения будет возрастать вместе с характеристикой появления ключа, которая обычно стремится к значению $H(K)$, хотя в некоторых случаях это и не так. В случае n -граммной подстановки, транспозиции, шифра Виженера, его вариантов и дробного шифра получаются строго идеальные системы для рассматриваемого простого языка и $H_E(M) \rightarrow H(K)$ при $N \rightarrow \infty$.

Идеальные секретные системы обладают следующими недостатками.

1. Система должна находиться в точном соответствии с языком. Это требует от создателя шифра глубокого изучения структуры языка. Кроме того, изменения статистической структуры или некоторый отбор сообщений из множества возможных сообщений, как в методе вероятных слов (слов, ожидаемых в данной частной криптограмме), делают возможным раскрытие системы.

2. Так как структура естественных языков крайне сложна, то для устранения избыточности требуются сложные преобразования. Следовательно, любое устройство, предназначенное для выполнения этой операции, необходимо должно быть очень сложным, по крайней мере в отношении хранения информации, так как можно ожидать, что потребуется «словарь», на порядок больший, чем обычные словари.

3. Вообще говоря, используемые отображения приводят к значительному разрастанию ошибок. Ошибка при передаче в одной букве приводит к ошибкам в области, объем которой сравним с порядком длительности статистических связей в исходном языке.

19. Дополнительные замечания о ненадежности и избыточности

Избыточность «нормативного английского языка» была взята равной 0,7 десятичных единиц на букву или 50%. Эта величина получена в предположении, что знаком пробела можно пренебречь. Она является приближенной величиной, подсчитанной на основе статистической структуры языка, учитывающей связи в пределах 8 букв, и при этом анализировался текст обычного типа: газетные статьи, литературные произведения и т. д. Можно указать один метод грубой оценки этой величины, которая представляет криптографически некоторый интерес.

Шифр бегущего ключа является системой типа Вернама, где ключ является не случайной последовательностью, а осмысленным текстом. Далее, известно, что шифры с бегущим ключом могут быть обычно решены единственным образом. Это говорит о том, что английский текст может быть сжат в два раза, откуда следует, что избыточность должна составлять не менее 50%. Однако эта цифра не может быть значительно увеличена по ряду причин (если

только не рассматривать далеко распространяющуюся «смысловую» структуру английского языка), приводящих к статистическим связям между далекими буквами.

Шифр бегущего ключа легко может быть усовершенствован таким образом, что система не может быть решена без ключа. Во-первых, если использовать вместо одного, скажем, четыре различных текста в качестве ключа, прибавляя все их к сообщению, то этим самым вводится объем ключа, достаточный для того, чтобы создать высокую положительную ненадежность. Во-вторых, в качестве ключа использовать, скажем, каждую десятую букву. Промежуточные буквы просто опускаются и дальше не используются. Это дает почти тот же самый эффект, что и предыдущий способ, так как такие разъединенные буквы почти независимы.

Тот факт, что для некоторого отрывка текста все гласные могут быть выброшены без существенных потерь смысла, наталкивает на простой способ улучшения почти всех секретных систем. Сначала вычеркнуть все гласные или по возможности больше букв сообщения, отсутствие которых не вызовет риска неоднозначного восстановления, а затем остаток зашифровать. Так как эта операция уменьшает избыточность в три или четыре раза, то отодвинется и точка единственности в три или четыре раза дальше. Одним из методов достижения идеальности секретной системы является использование в качестве части системы расшифровки знания английского языка шифровальщиком адресата.

20. Распределение ненадежности

Более полное описание секретной системы, применяемой к некоторому языку, чем то, которое дается характеристиками ненадежности, можно получить из рассмотрения *распределения ненадежности*. Для N перехваченных букв рассмотрим то множество криптограмм, для которых ненадежность и именно ненадежность при условии, что фиксированы эти конкретные криптограммы, а не средняя ненадежность $H_E(M)$, лежит в определенных пределах. Это дает плотность распределения

$$P(H_E(M), N) dH_E(M)$$

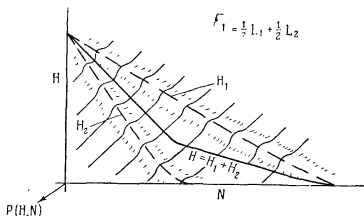
вероятности того, что для N букв ненадежность лежит между H и $H + dH$. Средняя ненадежность, рассмотренная выше, является средним значением этого распределения. Плотность $P(H_E(M), N)$ можно представлять себе отложенной вдоль третьей оси, перпендикулярной двум осям в плоскости $H_E(M), N$. Если язык является чистым, причем расстояния, на которых еще сказываются статистические связи, малы, а шифр также является чистым, то эта функция обычно образует подобие гребня, наивысшая точка которого

приблизительно соответствует $H_E(M)$. Это верно по крайней мере тогда, когда вблизи от наивысшей точки нет точки единственности. В этом случае, или в том случае, когда эти условия удовлетворены приближенно, кривая среднего значения распределения дает довольно полное представление о системе.

С другой стороны, если язык не является чистым, но составлен из нескольких чистых компонент

$$L = \sum p_i L_i,$$

имеющих различные характеристики ненадежности, то полное распределение будет состоять обычно из нескольких «гребней». Каждый из них будет соответствовать некоторому L_i , и они взвешиваются



Р и с. 11. Распределение ненадежности для смешанного языка $L = 1/2 L_1 + 1/2 L_2$.

в соответствии с весами p_i . Средняя характеристика ненадежности будет представлять собой линию, проходящую где-то посреди этих гребней и, возможно, не будет давать достаточно полную характеристику ситуации (см. рис. 11). Аналогичная картина получается и тогда, когда система не является чистой и составлена из нескольких систем с различными кривыми для функции.

В результате смешивания чистых языков, близких по статистической структуре, ширина гребней может увеличиться. Из-за этого вблизи точки единственности средняя ненадежность будет иметь тенденцию к возрастанию, так как ненадежность не может становиться отрицательной и расплывание происходит главным образом в положительном направлении. Поэтому можно ожидать, что в этой области вычисления, основанные на случайном шифре, должны давать до некоторой степени заниженный результат.

Часть III

ПРАКТИЧЕСКАЯ СЕКРЕТНОСТЬ

21. Рабочая характеристика

После того как объем перехваченного текста превзойдет точку единственности, обычно будет существовать единственное решение криптограммы. Задача криптографического анализа и состоит в выделении этого единственного решения, имеющего высокую вероятность. До того как достигнута точка единственности, задача криптографического анализа состоит в выделении всех возможных решений с большой (по сравнению с остальными решениями), вероятностью и в определении вероятностей этих решений.

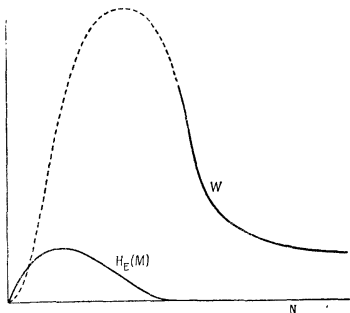
Хотя в принципе всегда можно найти эти решения (например, испытывая все возможные ключи), но для различных систем нужно будет затратить весьма различный объем работы. Средний объем работы, необходимой для определения ключа, если криптограмма имеет N букв, $W(N)$ измеренное, скажем, в человеко-часах, можно назвать рабочей характеристикой системы. Это среднее значение берется по всем сообщениям и всем ключам с соответствующими им вероятностями. Функция $W(N)$ измеряет количество «практической секретности» в данной системе.

Для простой подстановки, примененной к английскому языку, рабочая характеристика и характеристика ненадежности будут выглядеть примерно такими, как показано на рис. 12. Пунктирная часть кривой относится к области, где имеется несколько возможных решений и все они должны быть определены. Для сплошной части кривой после точки единственности существует, вообще говоря, только одно решение, но если необходимые данные имеются в минимальном количестве, то для нахождения этого решения нужно проделать большую работу. По мере увеличения количества данных объем необходимой работы быстро уменьшается, стремясь к некоторому асимптотическому значению, которое достигается, когда дополнительные данные уже не уменьшают объема работы.

По существу, такое же поведение кривых, как показано на рис. 12, можно ожидать для любого типа секретной системы, в которой ненадежность стремится к нулю. Однако масштаб количества требуемых человеко-часов будет сильно отличаться для разных типов шифров, даже тогда, когда кривые функции $H_E(M)$ почти совпадают. Например, шифр Виженера или составной шифр Виженера с тем же самым объемом ключа имели бы намного лучшие (т. е. более высокие) рабочие характеристики, чем простая подстановка. В хорошей (в практическом смысле) секретной системе график $W(N)$ остается достаточно высоким вплоть до того числа букв, которое намереваются передавать с помощью данного ключа,

что не дает противнику практической возможности найти решение или же задерживает нахождение решения до тех пор, пока содержащаяся в нем информация не устаревает.

В следующем разделе будут рассмотрены способы, которые позволяют удерживать функцию $W(N)$ на достаточно высоком



Р и с. 12. Типичная работа и характеристика ненадежности.

уровне даже в тех случаях, когда значение $H_E(K)$ практически равно нулю. Эта проблема по существу относится к типу «мини-максных» проблем, как всегда бывает в случае борьбы двух разумных соперников¹⁾. При создании хорошего шифра необходимо максимизировать минимальный объем работы, которую должен проделать противник для нахождения решения. Для этого недостаточно просто быть уверенным, что никаким из обычных криптографических методов нельзя раскрыть шифр. Надо быть уверенным, что этого нельзя добиться вообще никаким методом. И действительно, это оказалось слабой стороной многих секретных систем. Спроектированные так, чтобы быть недоступными для всех изве-

¹⁾ См. сноску на стр. 341. Ситуацию, складывающуюся между создателем шифра и противником, можно рассматривать как игру очень простой структуры: игра двух партнеров с нулевой суммой, с полной информацией и ровно с двумя «ходами». Создатель шифра в качестве своего «хода» выбирает некоторую систему. После этого противник информируется об этом выборе и выбирает метод анализа. «Цена» игры — это средняя работа, требуемая для решения криптограммы выбранным методом.

стных методов решения, они затем приводили к созданию новых криптографических методов, делавших возможным их раскрытие.

Проблема создания хорошего шифра является по существу проблемой нахождения наиболее сложных задач, удовлетворяющих определенным условиям. Это довольно нетривиальная ситуация, так как обычно в заданной области отыскивают простые и легко разрешимые задачи.

Как можно быть уверенным в том, что неидеальная система, которая необходимо имеет единственное решение для достаточно большого N , потребует большого объема работы для нахождения решения при любом методе анализа? Имеются два подхода к решению этого вопроса. 1) Можно изучить возможные методы решения, которыми пользуется противник, и попытаться описать их в достаточно общих терминах, с тем чтобы включить все методы, которые могли бы быть им использованы. Тогда наша система будет недоступной для этого «общего» метода решения. 2) Можно составить наш шифр таким образом, чтобы решение его было эквивалентно (или включало в себя) решению некоторой проблемы, про которую, известно, что для ее решения требуется большой объем работ. Так, если бы мы смогли показать, что нахождение решения некоторой секретной системы требует по меньшей мере столько же работы, сколько нужно для решения нераспадающейся системы уравнений с большим числом неизвестных, то можно получить некоторого рода нижнюю оценку для рабочей характеристики.

Следующие три раздела статьи посвящены этому вопросу. Трудно определить нужные понятия с точностью, достаточной для получения результатов в форме математических теорем, но думается, что выводы, полученные в форме общих принципов, остаются справедливыми.

22. Общие замечания о решении криптограмм

После того как количество перехваченного материала превзошло точку единственности, любая система может быть в принципе решена с помощью простого перебора всех возможных ключей до тех пор, пока не получится единственное решение, т. е. расшифрованное сообщение, «имеющее смысл» в исходном языке. Простое вычисление показывает, что этот метод решения (который можно назвать методом полной системы испытаний и ошибок) совершенно непрактичен, кроме тех случаев, когда ключ невероятно мал.

Предположим, например, что имеется ключ с $26!$ возможностями, (что составляет около $26,3$ десятичных единицы)¹⁾, т. е. того же самого объема, что и для простой подстановки в английском языке.

¹⁾ То есть $\log_{10} 26! \approx 26,3$. — Прим. ред.

Этот ключ при любом приемлемом способе измерения является малым ключом. Он может быть выписан на клочке бумаги, и его можно выучить в течение пяти минут. Его можно записать на 27-разрядном десятичном регистре или же на 88-ми разрядном двоичном регистре.

Предположим далее, давая противнику все возможные преимущества, что он сконструировал электронное устройство для испытания ключей со скоростью один ключ в одну *миксек* (скажем, с помощью автоматического выбора и проверки с помощью критерия значимости χ^2). Можно ожидать, что он выберет правильный ключ примерно после половины всех возможных испытаний, т. е. по прошествии примерно

$$2 \times 10^{26} / 2 \times 60^2 \times 24 \times 365 \times 10^6 = 3 \times 10^{12} \text{ лет.}$$

Другими словами, даже для малого ключа метод полной системы испытаний и ошибок не может быть использован на практике для решения криптограмм, исключая тот случай, когда ключ крайне мал, например в шифре Цезаря, где имеются только 26 возможностей (или 1,4 десятичных единицы). Метод испытаний и ошибок, обычно используемый в криптографии, несколько отличен от описанного выше, или же при использовании такого метода его дополняют с помощью других средств. Если бы кто-либо имел секретную систему, которая требует для раскрытия применения полной системы испытаний и ошибок, то он мог бы чувствовать себя в полной безопасности. Такая система получается, если исходные смысловые сообщения, состоящие все, скажем, из 1000 букв, выбираются случайным образом из множества всех 1000-буквенных последовательностей. Если к такому языку применить любой из простых шифров, то окажется, что нельзя было бы существенно улучшить метод полной системы испытаний и ошибок.

Методы, фактически используемые в криптографии, часто включают в себя метод испытаний и ошибок, но несколько отличного типа. Во-первых, испытания проводятся, начиная с более вероятных гипотез, и, во-вторых, каждое испытание относится к большой группе ключей, а не к единственному ключу. Так, пространство ключей можно разделить, скажем, на 10 подмножеств, содержащих примерно по одинаковому числу ключей. С помощью не более чем десяти испытаний находят правильное подмножество. Это подмножество делится на несколько вторичных подмножеств и процесс повторяется. Для нашего примера с объемом ключа $26! \approx 2 \cdot 10^{26}$ можно ожидать около $26 \cdot 5 = 130$ испытаний вместо 10^{26} при полной системе испытаний и ошибок. Этот результат может быть улучшен, если выбирать сначала наиболее вероятные из подмножеств. Если разделение производить на два подмножества (наилучший способ минимизировать число испытаний), то потребуются только

88 испытаний. В то время как метод полной системы испытаний и ошибок требует числа испытаний, равного по порядку числу ключей, испытания с разделением на подмножества требуют только число испытаний по порядку, равное объему ключа в битах.

Эти рассуждения остаются верными даже тогда, когда разные ключи имеют разные вероятности¹⁾. В этом случае соответствующий процесс минимизации среднего числа испытаний заключается в разделении пространства ключей на равновероятные подмножества. После того как отобрано нужное подмножество, оно снова делится на подмножества равной вероятности. Если этот процесс может быть продолжен, то среднее число испытаний для разделений на два подмножества будет равно

$$h = \frac{H(K)}{\log 2}.$$

Если каждая проверка имеет S возможных исходов, каждый из которых соответствует нахождению ключа в одном из S равновероятных подмножеств, то

$$h = \frac{H(K)}{\log S}$$

будет средним числом испытаний. Следует указать на интуитивное значение этих результатов. При проверке с разделением на два равновероятных подмножества каждое испытание дает один бит информации относительно ключа. Если подмножества имеют сильно отличающиеся вероятности, как при проверке единственного ключа в методе полной системы испытаний и ошибок, каждое испытание дает лишь малое количество информации. Так, для 26! равновероятных ключей проверка одного ключа дает только

$$-\left[\frac{26!-1}{26!} \log \frac{26!-1}{26!} + \frac{1}{26!} \log \frac{1}{26!} \right],$$

или приблизительно 10^{-25} бит информации. Разделение на S равновероятных подмножеств увеличивает количество информации, получаемой от одного испытания, до $\log S$, и среднее число испытаний равно полной информации, которая должна быть получена, т. е. $H(K)$, деленной на $\log S$.

Вопрос, затронутый здесь, имеет много общего с различными задачами о взвешивании монет, рассматривавшимися в последнее время. Типичным является следующий пример. Известно, что среди 27 монет одна фальшивая и она немного легче остальных. Требуется найти фальшивую монету с помощью ряда взвешиваний на аптекарских весах. Чему равно наименьшее число необходимых для

¹⁾ Более подробное обсуждение аналогичных вопросов можно найти в книге Я г л о м А. М. и Я г л о м И. М., Вероятность и информация, М., 1960.— Прим. ред.

этого взвешиваний? Правильный ответ — 3. Он получается, если сперва разделить все монеты на 3 группы по 9 в каждой. Две из этих групп сравниваются на весах. Три возможных результата определяют те 9 монет, среди которых находится фальшивая. Эти 9 монет снова делятся на 3 группы по 3 монеты и процесс продолжается. Множество монет соответствует множеству ключей, фальшивая монета — правильному ключу, а процесс взвешивания — испытанию или проверке. Первоначальная неопределенность равна 27 бит, и каждое испытание дает $\log_2 3$ бит информации. Таким образом, когда нет «диофантовых трудностей», то достаточно $\log_2 27 / \log_2 3$ испытаний.

Этот метод решения применим только тогда, когда пространство ключей может быть разделено на малое число подмножеств так, чтобы существовал простой способ определения подмножества, содержащего правильный ключ. Для того чтобы применить некоторый критерий совпадения и определить, подтверждено ли некоторое предположение, вовсе не нужно, чтобы это предположение относилось ко всему ключу — можно проверить предположение, относящееся только к части ключа (или предположение относительно того, лежит ли ключ в некоторой большой части пространства ключей). Другими словами, можно информацию о ключе получать бит за битом.

Возможность осуществления такого анализа обуславливает основную слабую сторону большинства секретных систем. Например, в простой подстановке некоторое предположение об одиночной букве можно проверить по ее частоте, разнообразию следующих за ней или предшествующих букв, по ее двоянным повторениям и т. п. При определении единственной буквы неопределенность пространства ключей (равная 26 десятичных единиц) может быть уменьшена на 1,4 десятичной единицы. Это справедливо для всех элементарных шифров. В шифре Виженера некоторое предположение относительно двух или трех букв ключа легко проверить следующим образом. Надо попытаться расшифровать другие части текста с помощью этого ключа и посмотреть, получится ли результат, имеющий смысл. Составной шифр Виженера намного лучше с этой точки зрения, если предположить, что его общий период больше, чем длина перехваченного текста. В этом случае при шифровании каждой буквы используется столько букв ключа, сколько имеется составляющих периодов. Хотя при этом используется только часть полного ключа, все же для применения удовлетворительной проверки требуется достаточно большое число букв.

Наш первый вывод состоит в том, что при разработке шифров с малым ключом надо стремиться к тому, чтобы при шифровании каждого малого элемента сообщения использовалась значительная часть ключа.

23. Статистические методы

Многие типы шифров могут быть раскрыты с помощью статистического анализа. Рассмотрим опять простую подстановку. Перехватив криптограмму, противник прежде всего должен произвести подсчет частот букв. Если криптограмма содержит, скажем, 200 букв, то без риска можно предположить, что лишь немногие из этих букв (если вообще такие найдутся) выйдут за пределы своих частотных групп, если группы получены с помощью разделения всех букв на 4 подмножества с четко отличающимися пределами частот. Логарифм числа ключей при этих ограничениях равен

$$\log 2! 9! 9! 6! = 14,28,$$

и, таким образом, простой подсчет частот уменьшает неопределенность ключа на 12 десятичных единиц — огромный выигрыш.

В общем случае статистический анализ выполняется следующим образом. По перехваченной криптограмме E вычисляется некоторая статистика. Эта статистика такова, что для всех осмысленных сообщений M она принимает значения, мало отличающиеся от S_h , величины, зависящей только от частного используемого ключа. Полученная таким образом величина служит для выделения тех возможных ключей, для которых значение S_h лежит в близкой окрестности наблюдаемого значения. Статистика, которая не зависит от K или изменяется в зависимости от M так же сильно, как и в зависимости от K , не может быть существенна для выделения некоторого подмножества ключей. Так, в шифрах транспозиции подсчет частот букв не дает никакой информации о K — для любого K эта статистика остается той же самой. Поэтому нельзя извлечь никакой пользы из подсчета частот для раскрытия шифров транспозиции.

Более точно данной статистике S можно приписать некоторую «разрешающую мощность». Для каждой величины S имеется условная ненадежность ключа $H_S(K)$ (ненадежность при фиксированном значении S) и это все, что известно относительно ключа. Взвешенное среднее этих величин

$$\sum P(S) H_S(K)$$

дает среднюю ненадежность ключа при известном S , где $P(S)$ является априорной вероятностью конкретного значения S . Разность объема ключа $H(K)$ и этой средней неопределенности измеряет «разрешающую мощность» статистики S .

В строго идеальном шифре все статистики данной криптограммы не зависят от частного используемого ключа. Это следует из свойства сохранения меры преобразованием $T_j T_k^{-1}$ в пространстве E или $T_j^{-1} T_k$ в пространстве M .

Имеются хорошие и плохие статистики, точно так же, как имеются хорошие и плохие методы испытаний и ошибок. Фактически проверка некоторой гипотезы методом испытаний и ошибок представляет собой некоторый тип статистики, и то, что было сказано выше относительно наилучших типов испытаний, верно и вообще. Хорошая статистика для решения системы должна обладать следующими свойствами.

1. Она должна просто вычисляться.
2. Она должна зависеть от ключа больше, чем от сообщения, если с ее помощью требуется находить ключ. Изменения по M не должны маскировать изменений по K .
3. Те значения статистики, которые могут быть «различены», несмотря на «размытость», создаваемую изменением по M , должны разделять пространство ключей на несколько подмножеств, вероятности которых сравнимы по величине, причем статистика будет характеризовать подмножество, в котором лежит правильный ключ. Статистика должна давать информацию о значительных объемах ключа, а не об объемах, составляющих малую долю общего числа бит.
4. Информация, даваемая статистикой, должна быть простой и удобной для использования. Таким образом, подмножества, на которые статистика разделяет пространство ключей, должны иметь простую структуру в пространстве ключей.

Подсчет частот букв для простой подстановки является примером очень хорошей статистики.

Можно предложить два метода (отличных от стремления приблизить систему к идеальной), которые будут мало доступны для статистического анализа. Их можно назвать методами «распыления» и «запутывания». В методе распыления статистическая структура сообщений M , которая приводит к избыточности в сообщениях, «распыляется» в статистику больших длин, т. е. в статистическую структуру, включающую длинные комбинации букв криптограммы. Тогда противник должен перехватить большой объем текста для восстановления этой структуры, так как она заметна лишь в блоках малой индивидуальной вероятности. Больше того, даже тогда, когда у противника достаточно материала, требуется гораздо больший объем вычислительной работы, так как избыточность распылена по большому числу индивидуальных статистик. Примером распыления статистик может служить «усреднение» сообщения $M = m_1 m_2 m_3 \dots$:

$$y_n = \sum_{i=1}^s m_{n+i} \pmod{26},$$

где складываются s последовательных букв сообщения для получения буквы y_n . Можно показать, что избыточность в последова-

Метод «запутывания» состоит в том, что соотношения между простыми статистиками в пространстве криптограмм и простыми подмножествами в пространстве ключей делаются весьма сложными и беспорядочными. Для случая простой подстановки легко описать ограничения на ключи, налагаемые частотами букв в криптограммах. Если эти ограничения очень сложны и беспорядочны, то противник, может быть, еще и может оценить, скажем, статистику S_1 , которая ограничивает некоторую область в пространстве ключей. Однако эта статистика выделяет некоторую весьма сложную область R_1 в пространстве ключей, возможно «свернутую» несколько раз, так что противнику трудно воспользоваться полученными результатами статистического анализа. Далее, вторая статистика S_2 ограничивает ключи областью R_2 и, следовательно, пересечением этих областей, но это не облегчает дела, так как трудно определить, что же именно представляет собой это пересечение.

$$f_1(k_1, \dots, k_p) = S_1,$$

• • • • •

$$f_n(k_1, \dots, k_p) = S_n,$$

Вывод, который можно сделать из приведенных рассуждений, состоит в том, что для улучшения секретной системы нужно применить тот или иной из рассмотренных методов или оба вместе.

24. Метод вероятных слов

Один из наиболее результативных приемов раскрытия шифров заключается в использовании вероятных слов. Вероятные слова — это слова или выражения, которые можно ожидать в частном сообщении вследствие того, что они характерны для данного источника сообщений. В качестве вероятных слов могут быть взяты просто общеупотребительные слова или отдельные слоги, которые встречаются в любом тексте на данном языке, такие, как *the*, *and*, *-tion*, *that* и т. п. для английского языка.

В общем случае метод вероятных слов применяется следующим образом. Предполагая, что некоторая часть криптограммы является определенным вероятным словом сообщения, находим весь ключ или часть ключа. Эта часть используется для расшифровки других частей криптограммы и служит критерием согласованности. Если другие части криптограммы становятся при этом понятными, то предположение подтверждается.

Имеется мало шифров классического типа, которые при небольшом ключе могут оставаться долго не раскрытыми методом анализа вероятных слов. Рассмотрение этого метода позволяет выработать некоторый способ проверки качества шифров, который можно было бы назвать «проверкой кислотой», как в ювелирном деле. Он применим только к шифрам с малым ключом (меньше 50 десятичных единиц), причем лишь тогда, когда шифры применяются к естественным языкам и в них не используется идеальный метод увеличения секретности. Метод «проверки кислотой» состоит в следующем. Решается вопрос: насколько трудно определить ключ или часть ключа, зная небольшую выборку из сообщения и соответствующую ей криптограмму? Любая система, в которой это можно легко сделать, не может быть очень трудно раскрываемой, так как шифровальщик всегда может применять метод вероятных слов в сочетании с методом испытаний и ошибок, до тех пор пока не будет получено правильное решение.

Условие, наложенное нами на объем ключа, делает количество испытаний малым, а условие относительно того, что система не должна быть идеальной, необходимо, так как идеальные системы автоматически дают методы проверки правильности решения. Наличие в криптограмме вероятных слов и фраз следует из предположения об использовании естественных языков.

Заметим, что требование, чтобы раскрытие шифра при этих условиях было трудным, само по себе не противоречит требованию простоты процессов шифрования и дешифрирования. Используя функциональные обозначения, запишем процесс шифрования следующим образом:

$$E = f(K, M),$$

а процесс дешифрирования —

$$M = g(K, E).$$

Обе эти операции могут быть простыми операциями над соответствующими аргументами, в то время как третья операция, записанная в виде

$$K = h(M, E),$$

может не быть простой.

Укажем, что при исследовании новых секретных систем один из наилучших методов их раскрытия состоит в рассмотрении того, как можно было бы определить ключ, если имеется достаточный объем сообщений M и соответствующая криптограмма E .

Идея метода запутывания может быть (и должна быть) использована для создания трудностей противнику, если он применит метод вероятных слов. Если шифровальщику противника даны (или он принял как данные) сообщение $M = m_1 m_2, \dots, m_s$ и криптограмма $E = e_1 \dots e_s$, то он может составить уравнения для различных элементов ключа $k_1 \dots k_r$ (а именно уравнение шифрования):

$$\begin{aligned} e_1 &= f_1(m_1, \dots, m_s; k_1, \dots, k_r), \\ e_2 &= f_2(m_1, \dots, m_s; k_1, \dots, k_r), \\ &\vdots \\ e_s &= f_s(m_1, \dots, m_s; k_1, \dots, k_r). \end{aligned}$$

Предположим, что в этих уравнениях известно все, за исключением k_i . Поэтому каждое из этих уравнений должно зависеть сложным образом от k_i и включать многие из них. В противном случае противник мог бы решить простые уравнения и, подставив результаты в более сложные уравнения, решить и их.

С точки зрения повышения запутанности желательно, чтобы функции f_i содержали несколько m_i , особенно если последние не являются смежными и, значит, менее коррелированы. Однако система при этом приобретает нежелательное свойство разрастания ошибок, так как тогда каждое e_i будет, вообще говоря, действовать на несколько m_i при дешифрировании, и ошибка будет распространена на все эти m_i .

Таким образом, для того, чтобы рабочая характеристика секретной системы была высокой, для получения любой буквы криптограммы должна быть использована как можно большая часть ключа с применением метода запутывания. Далее, желательна зависимость от нескольких несвязанных m_i , если можно допустить некоторое разрастание ошибок. Все три соображения, разобранные в этих разделах, подводят нас к рассмотрению метода «перемешивания».

25. Перемешивание

В некоторых разделах теории вероятностей очень ценным оказалось понятие перемешивания. Пусть имеется пространство с мерой (или вероятностное пространство) Ω и некоторое сохраняющее меру отображение F этого пространства в само себя, т. е. такое отображение, что мера отображенной области FR равна мере исходной области R . Отображение называется перемешиванием, если для любой функции, определенной на пространстве, и для любой области R интеграл от этой функции по области $F^n R$ стремится при $n \rightarrow \infty$ к интегралу от функции по всему пространству Ω , умноженному на объем области R . Это означает, что первоначальная область R перемешивается с равномерной плотностью по всему пространству, если F применяется большое число раз. В общем случае $F^n R$ становится областью, состоящей из большого числа тонких волокон, распределенных по всему пространству Ω . При увеличении n волокна становятся тоньше, а их плотность приближается к постоянной.

Перемешивание в этом точном смысле может осуществляться лишь в пространствах с бесконечным числом точек, так как в пространстве с конечным числом точек такое отображение должно быть периодическим. Однако можно понимать под перемешиванием, грубо говоря, такое отображение, которое распределяет любую разумно расположенную область довольно равномерно по всему пространству. Если до отображения область можно было бы описать в простых терминах, то после отображения потребовалось бы очень сложное описание.

В криптографии в качестве пространства Ω рассматриваются все возможные сообщения длины N , а в качестве области R — высоковероятные сообщения. Эта последняя группа сообщений имеет довольно простую статистическую структуру. Если применить перемешивание, то высоковероятные сообщения рассеялись бы равномерно по всему пространству.

Хорошее перемешивание часто получается с помощью повторных произведений двух простых не коммутирующих операций. Хопф¹⁾ показал, например, что масса теста может быть перемешана с помощью такой последовательности операций. Тесто сначала раскатывается тонким слоем, затем скатывается, затем опять раскатывается и т. д.

В хорошем перемешивании пространства с естественными координатами $X_1, \dots, X_s$ координата X_i переносится перемешиванием в точку X'_i с помощью соотношения

$$X'_i = f_i(X_1, \dots, X_s), \quad i = 1, 2, \dots, S,$$

¹⁾ H o p f E., On causality, statistics and probability, *Journal of Math. and Physics*, 13, (1934), 51—102.

причем функции f_i являются сложными, и зависят от всех переменных «чувствительным» образом. Небольшое изменение одного из переменных, скажем X_3 , значительно изменяет X'_i . Если X_3 проходят всю область возможных значений, то X'_i должны описывать длинный извилистый путь по пространству.

Можно придумать различные способы перемешивания, применимые к статистическим последовательностям такого типа, которые встречаются в естественных языках. Один из таких способов, который кажется достаточно хорошим, заключается в том, что сначала применяется транспозиция, а затем попеременно применяются подстановка и одно из простых линейных преобразований (например, сложение смежных букв по модулю 26). Таким образом, можно взять

$$F = LSLSLT,$$

где T — транспозиция, L — линейная операция, и S — подстановка.

26. Шифры типа T_hFS_j

Предположим, что F является хорошим перемешиванием, которое может быть применено к последовательности букв, а T_h и S_j — любые два семейства отображений, т. е. два простых шифра, которые могут и совпадать. Для конкретности будем считать, что оба они являются простыми подстановками.

Оказывается, что шифр TFS будет очень хорошей секретной системой с точки зрения рабочей характеристики. Во-первых, из наших рассуждений о статистических методах видно, что не существует простой статистики, дающей информацию о ключе. Любая важная статистика, полученная по криптограмме, должна быть очень сложной и чувствительной, так как в результате применения перемешивания F избыточность распыляется и запутывается. Метод вероятных слов также приводит к сложной системе уравнений (включающих все части ключа, если перемешивание хорошее), которые должны быть решены совместно.

Интересно отметить, что если отбросить шифр T , то остающаяся система будет подобна шифру S и поэтому не сильнее ее. Противник может просто аннулировать перемешивание с помощью операции F^{-1} , после чего решить полученную криптограмму. Если же отбросить шифр S , то остающаяся система будет намного сильнее, чем T (если перемешивание хорошее), но все же она еще будет несравнима с шифром TFS .

Этот принцип отделения простых шифров перемешивания может быть, конечно, продолжен. Например, можно было бы использовать секретную систему

$$T_hF_1S_jF_2R_1$$

с двумя перемешиваниями и тремя простыми шифрами. Эту систему можно упростить, используя один и тот же простой шифр и даже один и тот же ключ, так же как и одинаковое перемешивание. Это может сильно упростить механизацию применения таких систем.

Перемешивание, которое отделяет два (или более) ключа, действует для противника как своего рода фильтр — легко пронести через этот фильтр известный элемент, но неизвестный элемент (ключ) через него проходит не так легко.

Поставленная между двумя множествами неизвестных (множеством ключей S и множеством ключей T) операция перемешивания F «перепутывает» неизвестные вместе таким способом, что нахождение решения становится очень трудным.

Хотя системы, построенные по этому принципу, являются весьма надежными, они обладают одним существенным недостатком. Если перемешивание хорошее, то наблюдается сильное разрастание ошибок. Ошибка при передаче одной буквы влияет на несколько букв при дешифрировании.

27. Несовместимость требований к хорошим системам

Те пять критериев, предъявляемых к хорошим системам, которые были сформулированы в разд. 5, оказываются несовместимыми, если системы применяются к естественным языкам с их сложной статистической структурой. В случае искусственных языков с простой статистической структурой можно удовлетворить всем критериям одновременно с помощью шифров идеального типа. В естественных языках нужно идти на компромисс, учитывающий противоречивые требования, исходя из конкретных условий.

Если отбросить любой из этих пяти критериев, то оставшиеся четыре могут быть удовлетворены достаточно хорошо, как показывают следующие примеры.

1. Если не учитывать первое требование (количество секретности), то любая простая система (например, простая подстановка) будет удовлетворять остальным требованиям. В крайнем случае, когда это условие отброшено полностью, вообще не потребуется никакого шифра и можно посылать сообщение открытым текстом.

2. Если объем ключа не ограничен, то можно использовать систему Вернама.

3. Если ограничения не накладываются на степень сложности операций, то можно использовать крайне сложные типы приемов шифрования.

4. Если снять ограничение на разрастание числа ошибок, то весьма хорошей была бы система типа *TFS*, хотя она и несколько сложна.

5. Если допускается большое увеличение объема сообщения, то можно легко придумать различные системы, в которых «правильное» сообщение смешивается с многими «неправильными» сообщениями (дезинформация). Ключ определяет, какое из этих сообщений правильное.

Очень грубые соображения о несовместимости всех пяти критериев могут состоять в следующем.

Из условия 5 следует, что должны использоваться секретные системы, подобные тем, которые рассматривались в нашей статье, т. е. без большого числа пустых символов и т. п. Совершенные системы исключаются условием 2, а идеальные условиями 3 и 4. Итак, высокая секретность, требуемая условием 1, должна обуславливаться высокой рабочей характеристикой системы, а не высокой характеристикой ненадежности. Если ключ мал, система проста и ошибки не разрастаются, то метод вероятных слов, вообще говоря, решает систему довольно легко, так как в этом случае имеется сравнительно простая система уравнений для ключа.

Эти рассуждения слишком неопределенны для того, чтобы быть решающими, но их общий смысл кажется достаточно убедительным. Возможно, если бы удалось придать количественную определенность различным критериям, то можно было бы составить некоторые уравнения, связывающие количественно эти требования и дающие их оптимальные значения. Элементами, которые наиболее трудно поддаются численному измерению, являются сложность операций и сложность статистической структуры языка.

ПРИЛОЖЕНИЕ

Доказательство теоремы 3

Выберем произвольное сообщение M_1 и сгруппируем вместе все криптограммы, которые могут быть получены из M_1 с помощью любой операции шифрования T_i . Обозначим этот класс криптограмм C_i . Объединим в одну группу с M_1 все сообщения, которые выражаются в виде $T_i^{-1}T_jM_1$ и назовем этот класс C_1 . То же самое C_1 можно получить, если взять любое сообщение M из класса C_1 , так как

$$T_sT_j^{-1}T_iM_1 = T_iM_1.$$

При этом мы приходим и к тому же классу C_1 .

Выбирая некоторое сообщение M , не принадлежащее классу C_1 (если такое найдется), таким же способом построим классы C_2 и C_3 . Продолжая этот процесс, получим остаточные классы со

свойствами 1 и 2. Пусть M_1 и M_2 выбраны из класса C_1 . Предположим, что

$$M_2 = T_1 T_2^{-1} M_1.$$

Если криптограмма E_1 принадлежит классу C'_1 и может быть получена из сообщения M_1 с помощью отображений

$$E_1 = T_\alpha M_1 = T_\beta M_1 = \dots = T_\eta M_1,$$

то

$$E_1 = T_\alpha T_2^{-1} T_1 M_2 = T_\beta T_2^{-1} T_1 M_2 = \dots = T_\lambda M_2 = T_\mu M_2 = \dots.$$

Таким образом, каждое сообщение M_i из класса C_1 отображается в криптограмму E_1 с помощью одного и того же числа ключей. Аналогично каждая криптограмма E_i из класса C'_1 получается из сообщений, принадлежащих классу C_1 , с помощью того же самого числа ключей. Отсюда следует, что это число ключей является делителем полного числа ключей, и поэтому условия 3 и 4 выполнены.

Новейшие системы модуляции, такие, как ЧМ (частотная модуляция), ФИМ (фазово-импульсная модуляция) и КИМ (кодированно-импульсная модуляция), обладают интересным свойством, заключающимся в возможности взаимно замещать ширину полосы частот и отношение сигнал/шум; то есть они создают возможность передать ту же самую информацию передатчиком меньшей мощности, если использовать более широкую полосу частот. И наоборот, при использовании КИМ возможно уменьшить полосу частот за счет увеличения мощности сигнала. Открытие этих систем привело к пересмотру оснований теории связи. Ряд работ был посвящен этой области науки; среди них работы Габора, Винера, Таллера, Сулливана и автора.

Основные идеи теории связи не новы. Первые важные факты были установлены Найквистом и Хартли в 1920 г., а некоторые корни теории могут быть прослежены даже вплоть до работ физика девятнадцатого века Больцмана. В более новых исследованиях, однако, учитываются факторы, которые раньше игнорировались, в частности, теперь значительно глубже понимается действие шума в канале и важность статистических свойств сообщений, предназначенных к передаче.

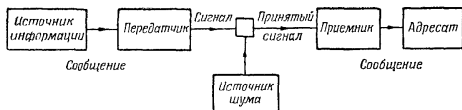
В данной статье основные стороны современной работы в этой области описываются с привлечением возможно простого математического аппарата. Это заставляет пожертвовать строгостью, поскольку рассматриваемый предмет является по существу математическим; более точное изложение читатель может найти в работах, указанных в конце статьи.

Тип системы связи, которая была исследована наиболее тщательно, показан на рис. 1. Система состоит из источника информации, который вырабатывает первичную информацию, или сообщения, предназначенные для передачи; передатчика, который кодирует

¹⁾ Shannon C., Recent development in communication theory, *Electronics*, April (1950), 80.

или модулирует эту информацию подходящим для канала способом; и канала, по которому закодированная информация, или сигнал, передается к пункту приема. Во время передачи сигнал может быть искажен шумом — на схеме указан источник шума. Принятый сигнал идет к приемнику, который декодирует или демодулирует его, чтобы восстановить первоначальное сообщение, а затем к пункту назначения информации.

Из дальнейшего будет видно, что эта система имеет достаточно общий вид для того, чтобы охватить большинство проблем связи,



Р и с. 1. Система связи, представленная схематически, в грубом приближении аналогична системе транспортировки грузов.

если соответствующим образом интерпретировать различные введенные элементы. В телевидении, например, источником информации является передаваемая сцена, сообщением — выходной сигнал телекамеры и сигналом — выходной сигнал передатчика.

Основная идея теории связи состоит в том, что с информацией можно обращаться почти так же, как с такими физическими величинами, как масса или энергия. Система, изображенная на рис. 1, грубо говоря, аналогична системе транспортировки; например, можно представить себе лесозавод, выпускающий бревна в определенном пункте, и систему транспортировки для перевозки бревен в другой пункт. В такой ситуации требуется решать два важных вопроса: определить величину скорости R (в кубических футах в секунду), с которой бревна выпускаются заводом, и величину пропускной способности (емкость) C (в кубических футах в секунду) системы транспортировки. Если R больше, чем C , то, конечно, невозможно перевезти всю готовую продукцию лесозавода. Если R меньше или равно C , то возможность перевозки находится в зависимости от того, могут ли бревна быть хорошо упакованы для транспортировки. Предположим, однако, что наш завод лесопильный. Тогда лесоматериал может быть распилен таким образом, чтобы использовать имеющуюся емкость транспортных средств на все 100%. Естественно, что в этом случае, прежде чем отослать эти пиломатериалы потребителю, их было бы надо направлять в плотничный цех приемного пункта для того, чтобы там придать им первоначальные габариты.

Если эта аналогия верна, то было бы можно выбрать меру R в подходящих единицах, показывающую, как много информации производится за секунду данным источником информации, и вторую меру C , которая определяет пропускную способность канала при передаче информации. Более того, было бы возможно при использовании подходящей кодирующей или модулирующей системы передавать информацию по каналу тогда и только тогда, когда скорость производства продукции R не больше, чем пропускная способность C . То, что это действительно возможно, является ключевым результатом современных исследований, и здесь будет вкратце указано, как это достигается.

Измерение информации

Прежде чем переходить к рассмотрению вопроса о том, как следует измерять информацию, нам необходимо объяснить точный смысл понятия информация с точки зрения инженера-связиста. Конечно, каждое подлежащее передаче сообщение имеет свое содержание. Оно, однако, совершенно несущественно в проблеме передачи информации. Передать ряд бессмысленных слогов так же трудно (в действительности даже более трудно), как и подлинный английский текст. Тому, кто хотя бы немного знаком с предметом этой статьи, будет ясно, что с точки зрения передачи важным свойством информации является то, что каждое частное сообщение выбирается из некоторого множества возможных сообщений. Передаче подлежит одно из частных сообщений, выбранных источником информации. Первоначальное сообщение может быть восстановлено в пункте приема в том и только в том случае, когда передается именно такое однозначным образом выбранное сообщение. Таким образом, информация в нашем смысле должна находиться в связи с понятием выбора из множества возможных исходов.

Простейшим типом выбора является выбор из двух возможностей, каждая с вероятностями $1/2$. Этот выбор осуществляется, например, когда кидают монету, которая с одинаковой вероятностью может упасть кверху гербом или решеткой. Удобно использовать количество информации, производимое таким выбором, в качестве основной единицы, называемой двоичной единицей или, короче, битом. Выбор, дающий один бит информации, может быть схематически изображен так, как это сделано на рис. 2, а. В точке b можно выбрать или верхнюю или нижнюю линию с вероятностями $1/2$ для каждой возможности. Если имеется N равновероятных возможностей, то количество информации равно $\log_2 N$. Это видно из рис. 2, б, где имеется восемь возможностей, каждая с вероятностью $1/8$. Выбор можно представлять себе происходящим в три этапа, каждый из которых дает информацию в один бит. Первый

бит соответствует выбору первых четырех или вторых четырех из восьми возможностей; второй бит соответствует первой или второй паре, выбранной из четырех возможностей, определенных первым выбором, а последний бит определяет первый или второй член этой пары. Ниже будет показано, что число требуемых битов равно $\log_2 N$, в нашем случае $\log_2 8$, т. е. 3.

Если вероятности не равны, то соответствующая формула несколько более сложна. Один простой случай показан на рис. 2, в. Здесь имеются четыре возможных выбора с вероятностями $1/2, 1/4, 1/8, 1/8$

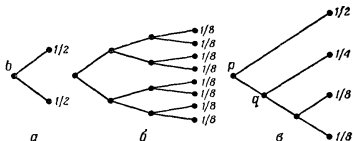


Рис. 2. Схематическое изображение равных и неравных вероятностей. Выбор, оцениваемый в один бит, можно сравнить с выпадением герба или решетки при бросании монеты.

и $1/8$. Этот выбор может быть разбит в последовательность двоичных выборов, как изображено на рис. 2, в. Произведенная информация задается числом $1 + \frac{1}{2} + \frac{1}{4}$, где 1 возникает из первого выбора (в точке p), который проводится всегда; $\frac{1}{2}$ — из выбора в точке q , который проводится лишь в половине случаев (когда в точке p выбрана нижняя линия), и т. д. Вообще, при подобном разложении, когда отдельные результаты выбора имеют вероятности $p_1, p_2, \dots, p_n$, информация задается формулой

$$H = -(p_1 \log_2 p_1 + p_2 \log_2 p_2 + \dots + p_n \log_2 p_n)^{1)}.$$
 (1)

Эта формула, таким образом, дает количество информации, произведенное одним выбором. Источник информации производит сообщение, которое состоит из последовательности, получающейся в результате выбора, например, букв печатного текста или отдельных слов, или звуков речи. В этих случаях количество информации, произведенное в секунду или даваемое одним символом, может быть вычислено с помощью соотношения (1). Интересно, что если статистическую структуру английского текста принимать во внимание лишь на расстояниях, не превышающих длину слова, то ско-

¹⁾ В литературе последних лет число H называют обычно энтропией.—
Прим. ред.

рость создания информации для печатного английского текста будет равна приблизительно двум битам на букву. Учет более далеких смысловых связей может значительно понизить эту величину.

Кодирование информации

Важность введения количественной меры информации H состоит в том, что она определяет возможную экономию времени, затрачиваемого на передачу, экономию, достижимую при соответствующем выборе системы кодирования, учитывающей статистические свойства источника сообщений. Для иллюстрации рассмотрим язык, в котором имеются лишь четыре буквы: А, В, С и D; пусть эти буквы имеют вероятности $1/2$, $1/4$, $1/8$ и $1/8$, как и на рис. 2, в.

В длинном тексте на этом языке А будет занимать половину всего текста, В — одну четверть и т. д. Предположим, что этот язык должен кодироваться двоичными знаками 0 и 1. Это может значить, что требуется передать текст посредством импульсной системы с двумя типами импульсов. Самый простой код таков: А—00, В—01, С—10, D—11. Этот код требует два двоичных знака на букву сообщения. При учете статистической природы текста можно сконструировать следующий, более хороший код: А—0, В—10, С—110, D—111. Легко проверить, что здесь исходное сообщение может быть восстановлено декодированием. Далее, число использованных двоичных знаков в среднем уменьшится. Действительно, оно вычисляется следующим образом:

$$\frac{1}{2} (1) + \frac{1}{4} (2) + \frac{1}{8} (3) + \frac{1}{8} (3) = 1 \frac{3}{4},$$

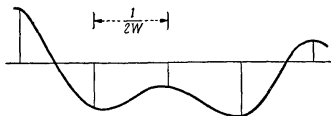
где первый член относится к букве А, которая встречается в половине всех случаев и которой соответствует один двоичный знак и т. д. Заметим, что $1\frac{3}{4}$ есть в точности значение H , вычисленное для рис. 2, в.

Результат, который был проверен для этого специального случая, выполняется всегда. Если на букву приходится в сообщении H битов, то возможно закодировать это сообщение с использованием в среднем только H двоичных знаков на букву текста. Не существует метода кодирования, который использует меньшее число двоичных знаков.

Пропускная способность канала

Теперь рассмотрим задачу определения пропускной способности C канала. Скорость работы источника информации измерялась в битах в секунду, и естественно измерять величину C в тех же единицах. Тогда возникает вопрос: «Какое максимальное число битов может быть передано в секунду по данному каналу?».

В некоторых случаях ответ прост. В канале телетайпа имеются 32 возможных символа. Каждый символ дает, следовательно, 5 битов



Р и с. 3. Сигналы, представленные функцией времени и лежащие в полосе частот W кГц, могут быть определены значениями последовательности выборок, отстоящих друг от друга на $1/(2W)$ секунд.

при условии, что возможные символы используются с равной вероятностью. Если можно передать n символов в секунду и уровень шума недостаточно высок для того, чтобы внести ошибку во время передачи, то можно передать $5n$ битов в секунду.

Предположим теперь, что канал определяется следующим образом. В качестве сигнала можно использовать любую функцию времени $f(t)$, которая лежит в некоторой полосе частот шириною в W герц. Известно, что функция этого типа может быть определена заданием ее значений в последовательности моментов выбора, отстоящих, как показано на рис. 3, друг от друга во времени на $1/(2W)$ секунд. Следовательно, можно сказать, такая функция имеет $2W$ степеней свободы (или измерений) в секунду.

Если шума нет, то в таком канале можно различить бесконечное число различных амплитудных уровней для каждого из выбранных моментов времени. Следовательно, в принципе можно было бы передать бесконечное число битов информации в секунду и пропускная способность C была бы бесконечной.

Если не ограничивать мощность передатчика, то пропускная способность будет бесконечной даже в случае наличия шума, так как все еще будет возможно отличать в каждый выбранный момент неограниченное число разных амплитудных уровней. Только в случае, когда имеется шум и мощность передатчика некоторым способом ограничена, получается конечная пропускная способность. Пропускная способность зависит, конечно, от статистической структуры шума так же, как и от того, каким образом ограничена мощность.

Простейшим типом шума является белый тепловой шум или шум сопротивления. В этом случае распределение вероятностей амплитуд является гауссовским, а спектр равномерен в указанной полосе частот и может быть положен равным нулю вне этой полосы. Этот тип шума полностью определяется заданием его средне-квад-

ратичной амплитуды N , которая равна мощности; ее следует измерять в стандартных единицах измерения сопротивления.

Простейший способ ограничить мощность передатчика состоит в том, чтобы предположить, что средняя мощность, создаваемая передатчиком (или, более точно, средне-квадратичная амплитуда сигнала), не больше чем P . Если определить наш канал этими тремя параметрами W , P и N , то пропускная способность C может быть вычислена. Она равна

$$C = W \log_2 \frac{P+N}{N} \quad (2)$$

битов в секунду. Легко увидеть, что эта формула приближенно верна, когда P/N велико. Полученный сигнал будет иметь мощность $P + N$ и возможно различать порядка $\sqrt{(P + N)/N}$ различных амплитуд в каждом сигнале. Причина этого состоит в том, что диапазон амплитуд полученного сигнала пропорционален $\sqrt{P + N}$, в то время как шум вносит неопределенность, пропорциональную $\sqrt{N}$. Количество информации, которое может быть передано одним мгновенным значением сигнала, будет, следовательно, равно $\frac{1}{2} \log_2 [(P + N)/N]$. Так как всего имеется $2W$ независимых моментов выбора значения сигнала в секунду, то пропускная способность дается формулой (2). Эта формула имеет значительно более глубокий и точный смысл, чем тот, на который указывают приведенные выше рассуждения. Действительно, можно показать, что при соответствующем выборе наших сигнальных функций можно передать $W \log_2 [(P + N)/N]$ битов в секунду со сколь угодно малой частотой ошибок. Невозможно ввести передачу с любой большей скоростью и с произвольно малой частотой ошибок. Это значит, что, несмотря на наличие шума, пропускная способность является четко определяемой величиной.

Формула для вычисления C применима для всех значений P/N . Даже в случае, когда P/N очень мало, а средняя мощность шума много больше средней мощности передатчика, возможно вести передачу со скоростью $W \log_2 [(P + N)/N]$ и с произвольно малой частотой ошибок. В этом случае $\log_2 (1 + P/N)$ очень близок к $(P/N) \log_2 e$ или к $1,443 P/N$, и приближенно $C = 1,443 P/N$.

Следует подчеркнуть, что передавать информацию по каналу со скоростью C возможно только при соответствующем кодировании информации. Вообще же говоря, скорость C не может быть в действительности достигнута; к ней можно только приближаться в пределе, используя все более и более сложные способы кодирования и при все большем запаздывании в работе передатчика и приемника. При белом шуме наилучшее кодирование таково, что переданные сигналы имеют структуру шума сопротивления мощности P .

Идеальные и практические системы

На рис. 4 изображена функция $C/W = \log(1 + P/N)$, рассматриваемая как функция отношения P/N , измеренного в децибеллах. Она представляет собой, следовательно, пропускную способность канала с белым шумом на единицу полосы частот. Кружки и точки соответствуют системам КИМ и ФИМ, используемым для

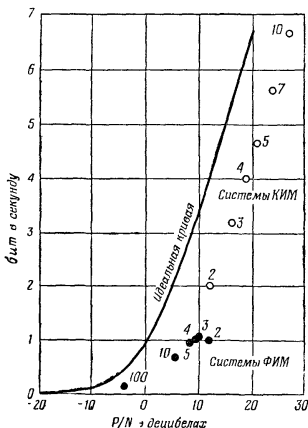
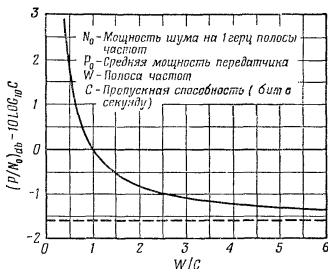


Рис. 4. Пропускная способность канала на единицу полосы частот как функция отношения сигнал/шум для двоичного канала.

посылки последовательности двоичных символов и приспособленных к тому, чтобы давать примерно одну ошибку на 10^5 битов. В случае КИМ число, стоящее около точки, представляет собой число амплитудных уровней; например, 3 означает систему КИМ с тремя уровнями. Во всех случаях используются положительные и отрицательные амплитуды. Системы ФИМ квантованы дискретным множеством возможных значений импульса, отстоящих друг от друга на $1/(2W)$; число, стоящее рядом с точкой, равно числу возможных значений импульса.

Совокупность точек расположена на кривой того же вида, что и в идеальном случае, но смещенной по горизонтали примерно на 8 децибелл. Это значит, что при более сложных способах кодирования и модуляции мог бы быть достигнут выигрыш в мощности в 8 децибелл по сравнению с указанными системами.



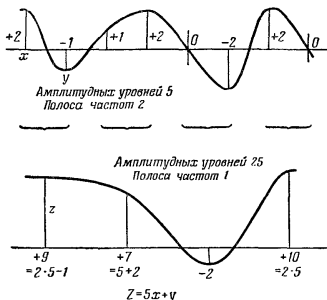
Р и с. 5. Если отношение сигнал/шум велико, уменьшение вдвое полосы частот увеличивает вдвое отношение сигнал/шум для данной пропускной способности канала.

К несчастью, при попытках достичь идеала передатчик и приемник с необходимостью становятся все более и более сложными и возрастают задержки. В силу этих причин наступает момент, когда устанавливается некоторое экономическое равновесие между различными факторами. Возможно тем не менее, что даже в настоящее время были бы оправданы более сложные системы.

Курьезный факт, иллюстрирующий общее мизантропическое свойство природы, состоит в том, что в обоих экстремумах P/N (когда точки лежат вне практически интересной области) ряды точек на рис. 4 ближе подходят к идеальной кривой.

Соотношение $C/W \log(1 + P/N)$ можно рассматривать как зависимость между параметрами W и P/N . Считая пропускную способность фиксированной, можно уменьшить ширину полосы частот W при условии, что отношение P/N будет достаточно увеличено и, наоборот, увеличение полосы приводит к уменьшению отношения сигнал/шум в канале. Требуемое значение P/N в децибеллах показано на рис. 5 как функция полосы частот W . Здесь предполагается, что когда увеличивается полоса частот W , мощность шума N возрастает пропорционально: $N = W N_0$, где N_0

есть мощность шума, приходящаяся на 1 герц полосы. Заметим, что если P/N велико, то сокращение полосы частот очень невыгодно с точки зрения мощности. Уменьшение полосы в два раза, говоря грубо, удваивает требуемое отношение сигнал/шум в децибеллах.



Р и с. 6. Графическое изображение системы связи, в которой за счет повышения мощности передатчика полоса частот сохраняется неизменной.

Один из методов замещения ширины полосы частот отношением сигнал/шум показан на рис. 6. Верхняя кривая представляет собой сигнал, полоса частот которого такова, что она может быть однозначно определена при задании ее значений в указанные моменты выбора. В каждый момент выбора сигнал имеет 5 амплитудных уровней. Нижняя кривая получена, как показано, посредством комбинирования попарно значений в моменты выбора для первой кривой. Теперь имеется 25 амплитудных уровней, которые должны быть различены, но зато на этой кривой моменты выбора встречаются вдвое реже: следовательно, полоса частот уменьшена вдвое за счет удвоения отношения сигнал/шум. Операция, обратная этой, удваивает полосу, но уменьшает требуемое отношение сигнал/шум.

Резюмируя изложенное, следует сказать, что в системе, такой, как телевидение или передача речи, имеются три существенно различных способа, посредством которых может быть уменьшена ширина полосы частот. Первый из них явным образом замещает ширину полосы частот отношением сигнал/шум, используя только что указанный способ. Второй метод основан на использовании статистических связей, имеющих в сообщении. Этот метод основан

вается на частных свойствах источника информации и может рассматриваться как один из приемов согласования источника с каналом. Наконец, могут быть использованы особенности получателя информации. Так, при передаче речи ухо сравнительно нечувствительно к искажению фазы. Следовательно, информация о фазе не так важна, как информация об амплитуде и ее не нужно посылать с такой же точностью. Этот факт может быть использован для уменьшения полосы частот, и действительно, часть выигрыша, достигаемого применением вокодера, обусловлена этим эффектом. Вообще, использование частных особенностей восприятия информации получателем требует соответствующего согласования канала с требованиями данного получателя.

Многие современные системы связи в высшей степени неэффективны в отношении использования статистических свойств источника информации. Для иллюстрации этого рассмотрим систему передачи английской речи (не музыки и не других звуков). Требования, предъявляемые к качеству репродуцируемой речи, состоят только в возможности уловить ее смысл. Особенности произношения, интонация и другие аналогичные характеристики речи говорящего могут быть потеряны в процессе передачи. В этом случае можно было бы, по крайней мере в принципе, вести передачу по следующей схеме. При передатчике устанавливается устройство, которое печатает английский текст, соответствующий произносимым словам. Текст может быть закодирован в двоичных единицах, причем в среднем используется не более двух двоичных единиц на букву или девять на слово. Считая 100 слов в минуту приемлемой скоростью, получаем в случае, когда понятность является единственным требованием к точности передачи в качестве оценки скорости создания информации на английском языке, цифру 15 битов в секунду. Согласно рис. 4, эта информация могла бы быть передана по каналу с отношением сигнал/шум в 20 децибелл и с шириной полосы частот только в 2,3 герца!

ЛИТЕРАТУРА

- Hartley R.V.L., Transmission of Information, *B.S.T.J.*, July (1928), 535.
Shannon C. E., A mathematical theory of communication, *B. S. T. J.*, July (1948), 379 and Oct. (1948), (русский перевод см. стр. 244 данного сборника.—Прим. ред.)
Shannon C. E., Communication in the presence of noise, *Proc. IRE*, Jan. (1949), (русский перевод см. стр. 433 данного сборника.—Прим. ред.)
Feldman C. B., Bennett W. R., Bandwidth and transmission performance, *B. S. T. J.*, July (1949), 490.
Fink D. G., Bandwidth vs noise in communication systems, *Electronics*, Jan. (1948).

Краткое содержание

В новейших работах [1—6] описываются опыты по передаче речи при помощи кодово-импульсной модуляции (КИМ). В этой статье в общем виде показаны преимущества КИМ и установлены различия в возможностях КИМ и других широкополосных систем, как, например, ЧМ с большим индексом. Ставится цель — по возможности просто разъяснить ряд вопросов, не входя в детали и не затрагивая некоторые вопросы, возникающие при конструировании соответствующих устройств.

1. КИМ и ее свойства

В систему КИМ входит ряд важных элементов. Эти элементы и роль, которую они играют, будут описаны в данном разделе.

Отсчитывание

Назначением системы передачи является, вообще говоря, воспроизведение на выходе системы некоторой функции времени, заданной на входе. Во всякой реальной системе приходится иметь дело только с определенным классом функций на ее входе, а именно с функциями с ограниченным спектром. Сигнал, не содержащий частот выше W_0 , не может представляться бесконечным числом независимых значений в секунду. Он может в действительности представляться в точности $2W_0$ независимыми значениями в секунду, и совокупность значений, отстоящих друг от друга во времени на τ_0 секунд, где $\tau_0 = 1/(2W_0)$ определяет сигнал полностью. Простое доказательство этого дано в приложении I. Итак, чтобы передать сигнал с ограниченным спектром длительностью T , не требуется передавать полностью всю непрерывную функцию времени. Достаточно передать конечную совокупность $2W_0 T$ независимых величин, получаемых путем отсчета мгновенных значений сигнала с постоянной скоростью $2W_0$ отсчетов в секунду.

¹⁾ Oliver B., Pierce J., Shannon C., The philosophy of P. C. M., *Proc. IRE*, 36, 11 (1948), 1324.

Если читателю покажется удивительным, что $2W_0T$ отдельных данных описывают полностью непрерывную функцию на интервале T , то следует напомнить, что для такого описания достаточно $2W_0T$ коэффициентов ряда Фурье, которым может быть представлена наша функция, принимая во внимание, что она не содержит частот выше W_0 .

Восстановление

Перейдем теперь к приемному концу системы и предположим, что отдельные значения, представляющие сигнал, поступают в надлежащей временной последовательности и могут быть использованы с неизменной скоростью $2W_0$. Чтобы восстановить сигнал, необходимо просто создать пропорциональный каждому данному значению импульс и пропустить полученную последовательность регулярно следующих друг за другом импульсов через идеальный фильтр нижних частот с граничной частотой W_0 . На выходе фильтра получим в точности (не считая общего запаздывания и возможно коэффициента пропорциональности) первоначальный сигнал. Так как откликом идеального фильтра нижних частот на короткий импульс является импульс вида $\sin x/x$ и так как общий эффект на выходе есть сумма откликов на все воздействия на входе, то описанный метод восстановления сигнала является просто физическим осуществлением метода, указанного в приложении I.

Итак, в идеальном случае имелось бы совершенное воспроизведение сигнала, если бы передавалась информация, дающая нам в точности мгновенные значения сигнала через интервалы $1/(2W_0)$, во времени.

Квантование

Невозможно, конечно, передать *точное* мгновенное значение. Мгновенное значение сигнала передается часто как амплитуда импульса или как его положение во времени. Помехи, искажения и взаимное влияние импульсов изменяют высоту и положение и вносят ошибку в полученную информацию о величине отсчета. Обычно ошибка возрастает по мере того, как сигнал усиливается последовательными повторителями, и накопление шума кладет предел расстоянию, на которое сигнал может быть передан даже при наличии достаточного усиления.

Возможно, впрочем, ограничиться только некоторыми дискретными уровнями амплитуды или положения передаваемого импульса. Тогда при взятии отсчета должен передаваться уровень, ближайший к истинному. После приема и усиления получим уровень, немного отличающийся от какого-либо из установленных. Если помехи

и искажения не слишком велики, можно с уверенностью сказать, какой именно уровень должен был бы иметь сигнал. После этого сигнал может быть заново сформирован или может быть создан новый сигнал, имеющий снова первоначально переданный уровень.

Представление сигнала некоторыми установленными дискретными уровнями называется квантованием. Оно неизбежно вносит ошибку в определение величины отсчетов, порождая шум квантования. Но если сигнал находится в квантованном состоянии, он может передаваться на любое расстояние без дальнейшей потери качества, если только добавочный шум в сигнале, принимаемом каждым повторителем, не настолько велик, чтобы нельзя было распознать правильный уровень каждого данного сигнала. Квантование сокращает наш «алфавит». Если принятый сигнал лежит между a и b и ближе, скажем к b , то можно считать, что передано b . Если шум достаточно мал, ошибки не будет.

Кодирование

Квантованный отсчет может быть передан в виде одиночного импульса, имеющего либо некоторую возможную дискретную амплитуду, либо некоторое дискретное положение по отношению к исходному. Однако если требуется большое количество различных амплитуд, например 100, то затруднительно сделать схемы, способные отличить их друг от друга. С другой стороны, очень легко осуществить схему, которая отличала бы наличие импульса от его отсутствия. Положим, далее, что несколько импульсов образуют кодовую группу для описания значения отдельного отсчета. Каждый импульс может быть налицо (1) или отсутствовать (0). Если взять три импульса, то можно составить комбинации, приведенные в следующей таблице:

Таблица I

Представляемое значение	0	1	2	3	4	5	6	7
Код	000	001	010	011	100	101	110	111

Кодовые обозначения представляют собой по существу числа, записанные в двоичной системе. В этой системе разряды соответствуют числам 1, 2, 4, 8, т. е. единица в первом разряде изображает 1, единица во втором разряде изображает 2, единица в третьем разряде изображает 4 и т. д. Кодовыми группами из n двоичных («да — нет») импульсов можно представить 2^n значений. Например, семь импульсов дают 128 уровней отсчета.

Возможно, конечно, кодировать мгновенное значение при помощи некоторого количества импульсов, которым присвоены величины 0, 1, 2 (основание 3; троичный код) или 0, 1, 2, 3 (основание 4; четверичный код) и т. д., вместо импульсов с присвоенными величинами 0 1 (основание 2; двоичный код). Если бы импульсу были присвоены десять уровней, то каждый импульс в кодовой группе представлял бы собой просто цифру в обычном десятичном числе, выражающем мгновенное значение сигнала. Если n есть число импульсов в группе, а b — основание, то число квантованных уровней, могущих быть представленными кодом, есть b^n .

Декодирование

Для того чтобы декодировать кодовую группу вышеописанного типа, нужно создать импульс, представляющий собой линейную сумму всех импульсов группы, умноженных на соответствующий вес разряда ($1, b, b^2, b^3, \dots$) в коде. Это может быть сделано многими способами. Простейший, возможно, состоит в том, что кодовая группа передается в обратном порядке, т. е. сначала «единицы», а наиболее высокий разряд последним. Импульсы накапливаются в форме заряда на комбинации RC (сопротивление—емкость) с такой постоянной времени, чтобы заряд убывал в b раз за время между импульсами. После получения последнего импульса заряд (напряжение) отсчитывается.

Полная система КИМ

Система КИМ включает все описанные выше процессы. Входной сигнал ограничивается по спектру так, чтобы исключить все частоты выше W_0 . После этого с сигнала снимаются отсчеты со скоростью $2W_0$. Отсчеты затем квантуются и кодируются. Так как возможны только дискретные кодовые группы, то выбор ближайшей кодовой группы автоматически квантует отсчет и в некоторых типах устройств нет необходимости в квантовании как в отдельной предварительной операции. Кодовые группы затем передаются либо в виде временной последовательности импульсов (временное разделение) по одному и тому же каналу, либо путем частотного разделения, либо по раздельным каналам. Кодовые группы регенерируются (восстанавливаются) с требуемым интервалом. В приемнике регенерированные кодовые группы декодируются и образуют последовательность импульсов, пропорциональных первоначальным отсчетам (не считая квантования); эти импульсы пропускаются через фильтр нижних частот с граничной частотой W_0 для восстановления первоначального сигнала.

2. Условия передачи для КИМ

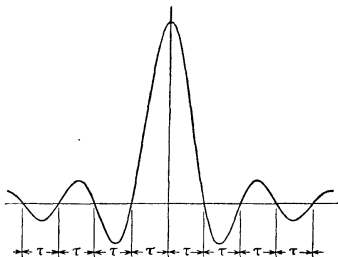
Рассмотрим условия, имеющие место в идеале в канале, предназначенном для передачи кодированных КИМ сигналов. То есть исключим физически невозможные приборы, но будем предполагать такие элементы, как, например, фильтры, селекторы и т. п., идеальными.

Ширина полосы

Если канал имеет ширину полосы W герц, то по нему можно передать $2W$ независимых импульсов в секунду. Это можно показать очень просто. Пусть импульсы происходят (или не происходят) в моменты $t = 0, \tau, 2\tau, \dots, m\tau$, где $\tau = 1/2W$, и пусть каждый принятый импульс имеет форму

$$V = V_0 \frac{\sin \frac{\pi}{\tau} (t - m\tau)}{\frac{\pi}{\tau} (t - m\tau)}, \quad (1)$$

представленную графиком рис. 1. Как видим, импульс, ось симметрии которого приходится на момент $m\tau$, будет равен нулю при $t = k\tau$,



Р и с. 1. Импульс вида $v_0 \frac{\sin(\pi t/\tau)}{\pi t/\tau}$.

где $k \neq m$. Таким образом, если берется отсчет последовательности импульсов при $t = m\tau$, то имеется только импульс, соответствующий этому моменту, но ни одного из остальных.

Далее, импульс, выражаемый формулой (1), не содержит частот выше W . Импульс такой формы возникает на выходе идеального фильтра нижних частот при воздействии на его вход весьма короткого импульса.

Чтобы передать сигнал с шириной полосы W_0 при помощи КИМ, надо посылать $2W_0$ кодовых групп в секунду, и каждая группа состоит, скажем, из n импульсов. Следовательно, необходимо передавать $2nW_0$ импульсов в секунду, а это требует полосы $W = nW_0$. Импульсы могут передаваться в форме временной последовательности по одному каналу или путем частотного разделения. В обоих случаях ширина полосы будет той же самой. Конечно, если в случае частотного разделения применяется передача с обеими боковыми полосами или если при временном разделении сигнал передается в форме радиочастотных импульсов с обеими боковыми полосами, полная ширина полосы будет $2nW_0$.

Коротко говоря, полоса частот, необходимая для КИМ, в идеальном случае в n раз больше той, которая нужна для непосредственной передачи сигнала; здесь n означает число импульсов в кодовой группе.

Пороговая мощность

Для уверенного обнаружения наличия или отсутствия импульса необходимо определенное отношение сигнала к помехе. Если мощность импульса слишком мала по сравнению с помехой, то и наилучший возможный приемник будет делать ошибки и указывать наличие импульса там, где его нет, и наоборот. Положим, что имеется идеальный приемник, т. е. такой, который делает наименьшее возможное количество ошибок. Если принимаемые импульсы имеют форму (1), а помеха представляет собой «белый» шум (т. е. шум с однородным спектром и гауссовским распределением, как, например, тепловой шум), идеальный прием может быть осуществлен путем пропускания сигнала через идеальный фильтр нижних частот с граничной частотой W (nW_0 в идеальном случае) и снятием отсчетов в моменты kt . Если при отсчете сигнал превосходит $V_0/2$, считаем, что импульс налицо, а если он меньше $V_0/2$, что импульс отсутствует. Результат будет ошибочным, если шум в данный момент превосходит $V_0/2$ в данном направлении. При гауссовском распределении вероятность такого события пропорциональна дополнительной функции ошибок¹⁾ от

$$\frac{V_0}{2\sigma} = \sqrt{\frac{P_s}{4N}},$$

где σ — средне-квадратичное значение шума, $P_s = V_0^2$ — «мощность» сигнала (импульса), $N = \sigma^2$ — мощность шума в полосе W . По мере возрастания мощности сигнала P_s эта функция убывает очень быстро, так что если P_s/N достаточно велико, чтобы сигнал

¹⁾ Дополнительная функция ошибок от x есть $\frac{1}{\sqrt{2\pi}} \int_x^{\infty} e^{-\lambda^2/2} d\lambda$.

был только различим, то небольшое увеличение делает передачу почти идеальной. О том, насколько быстро происходит это улучшение, можно составить себе представление из табл. II. Цифры последнего столбца относятся к частоте следования импульсов 10^5 в секунду.

Ясно, что имеется вполне определенный *порог* (около 20 дБ), ниже которого искажение значительно и выше которого оно пренебрежимо. Сравнивая эту цифру 20 дБ с 60—70 дБ, необходимых при прямой высококачественной АМ¹⁾ передаче речи, видно, что КИМ

Таблица II

Сигнал/шум $\frac{P_s}{N}$	Вероятность ошибки	Это составляет одну ошибку за
13,3	10^{-2}	10^{-3} сек
17,4	10^{-4}	10^{-1} сек
19,6	10^{-6}	10 сек
21,0	10^{-8}	20 мин
22,0	10^{-10}	1 день
23,0	10^{-12}	3 месяца

требует значительно меньшей мощности сигнала, даже учитывая n -кратное возрастание мощности шума из-за увеличения полосы.

В предыдущих рассуждениях предполагалось применение двоичной системы. В этой системе импульсы имеются в среднем в течение половины всего времени и средняя мощность сигнала будет $P_s/2$ (см. приложение II). Если применяется симметричная двойная система, в которой 1 передается как (+)-импульс, а 0 — как (—)-импульс, то размах изменения сигнала должен быть таким же, как в системе «да — нет» для данного уровня шума, и этот размах может быть достигнут при применении импульсов половинной амплитуды. Так как всегда имеется либо (+)-, либо (—)-импульс, мощность сигнала будет равна $P_s/4$.

Если применяются импульсы, имеющие b различных уровней амплитуды (т. е. система с основанием b), то необходим определенный интервал между соседними уровнями для обеспечения перекрытия шума. Обозначим этот интервал $K\sigma$, где K — постоянная (из предыдущего видно, что K имеет значение около 10). Полный диапазон амплитуд составляет $K\sigma(b-1)$. Мощность сигнала будет наименьшей, если диапазон амплитуд симметричен относительно нуля, т. е. лежит в пределах от $-\frac{K\sigma}{2}(b-1)$ до $+\frac{K\sigma}{2}(b-1)$. Сред-

¹⁾ Амплитудная модуляция. — Прим. ред.

ная мощность сигнала в предположении равновероятности всех уровней будет (см. приложение II) равна

$$S = K^2 \sigma^2 \frac{b^2 - 1}{12} = K^2 N \frac{b^2 - 1}{12}. \quad (2)$$

Нужно отметить, что требуемая мощность сигнала быстро возрастает с увеличением основания b .

Регенерация

В большинстве систем передачи шум и искажения от отдельных звеньев накапливаются. При данном качестве передачи системы в целом требования к каждому звену тем более суровы, чем длиннее система. Так, например, если цепь состоит из 100 звеньев, то мощность шума, добавляемая каждым звеном, может составлять лишь одну сотую от величины, допустимой в случае одного-единственного звена.

Сигнал в системе КИМ можно восстанавливать сколь угодно часто; поэтому амплитудные, фазовые и нелинейные искажения в одном звене, если они не слишком велики, не оказывают никакого влияния на восстановленный входной сигнал следующего звена. Если вследствие шума при наличии одного звена восстанавливается неправильно некоторая доля p всех импульсов, а $p \ll 1$, после m звеньев, доля эта возрастет приблизительно до mp . Однако для уменьшения p до значения $p' = \frac{P}{m}$ требуется, как было показано выше (раздел «Пороговая мощность»), лишь незначительное увеличение мощности в каждом звене. Следовательно, практически условия передачи для одного звена почти не зависят от общей длины системы. Значение этого обстоятельства едва ли можно переоценить.

3. Действие системы КИМ

Мы видим, что КИМ требует большей ширины полосы и меньшей мощности, чем прямая передача самого сигнала или простая амплитудная модуляция. В известном смысле произошел обмен ширины полосы на мощность. Выгодна ли эта замена? Достижимы ли большие отношения сигнал/шум в восстановленном сигнале? Насколько помехоустойчива система КИМ? Постараемся ответить на эти вопросы.

Пропускная способность канала

Хорошей мерой использования полосы является пропускная способность системы по сравнению с теоретическим пределом для канала с такой же шириной полосы и мощностью. Пропускная

способность системы может быть определена как число независимых символов или знаков, которые могут быть переданы без ошибок в единицу времени. Простейшим наиболее элементарным знаком является двоичная цифра, и удобно выражать пропускную способность числом C двоичных цифр в секунду, которое может пропустить канал. Шеннон и другие показали, что идеальная система имеет пропускную способность¹⁾

$$C = W \log_2 \left(1 + \frac{P}{N} \right), \quad (3)$$

где W — ширина полосы, P — средняя мощность сигнала, N — мощность белого шума. Два канала, имеющих одинаковые C , способны передавать одинаковое количество информации, если даже W , P и N различны.

В системе КИМ, работающей с порогом, таким, что частота ошибок пренебрежимо мала,

$$C = sm,$$

где s — частота отсчетов, равная $2W_0$, m — эквивалентное число двоичных цифр на кодовую группу. Если l означает число уровней квантования, то число двоичных цифр на кодовую группу есть $l = 2^m$, тогда как действительное число n цифр при основании b равно $l = b^n$.

Таким образом,

$$\begin{aligned} 2^m &= b^n, \\ m &= n \log_2 b \end{aligned}$$

и

$$C = sn \log_2 b.$$

Произведение sn выражает частоту следования импульсов, равную в идеальном случае удвоенной ширине полосы W . Поэтому

$$C = 2W \log_2 b = W \log_2 b^2.$$

Подставляя вместо b мощность, потребную при данном основании, из выражения (2) получаем

$$C = W \log_2 \left(1 + \frac{12S}{K^2 N} \right).$$

Сравнивая (4) и (3), видим, что они совпадают при $S = \frac{K^2}{12} P$. Другими словами, КИМ требует в $\frac{K^2}{12}$ (т. е. около 8) раз большей

¹⁾ Shannon C., A mathematical theory of communication, *BSTJ*, 27, July — October (1948). (См. стр. 273 настоящего сборника.—Прим. ред.)

мощности, чем теоретически необходимо для достижения заданной пропускной способности при данной ширине полосы.

Наиболее существенное замечание по поводу формулы (4) состоит в том, что формула правильна. Мощность и ширина полосы взаимно заменяются в логарифмическом отношении, и пропускная способность канала пропорциональна W^1). В большинстве широкополосных систем, улучшающих отношение сигнал/шум за счет расширения полосы, S пропорционально только $\log W$.

Отношение сигнал/шум

Система КИМ вносит помехи двоякого рода. Одна из них есть шум квантования, упомянутый выше в разделе «Квантование». Эта помеха вносится на передающем конце системы и больше нигде. Другая помеха есть шум ложных импульсов, происходящий вследствие ошибочной оценки значения импульса приемником или любым повторителем. Этот шум может возникать в любом участке системы и обладает свойством накапливания. Однако, как было показано выше, этот шум при увеличении превышения мощности сигнала над порогом убывает так быстро, что в любой реальной системе он может быть сделан пренебрежимо малым путем правильного конструирования. В результате отношение сигнал/шум в системе КИМ зависит только от шума квантования.

Если сигнал велик по сравнению с отдельным уровнем квантования, ошибки, вносимые квантованием в последовательные отсчеты, практически независимы. Наибольшая возможная ошибка составляет половину уровня в каждом направлении. Все значения ошибки вплоть до этого максимума равновероятны. Средне-квадратичное значение вносимой ошибки равно поэтому высоте уровня, умноженной на $1/(2\sqrt{3})$ (см. приложение II). Когда сигнал восстанавливается из декодированных отсчетов (содержащих эту ошибку квантования), то получается первоначальный сигнал плюс шум с однородным спектром вплоть до W_0 и со средне-квадратичной амплитудой, равной высоте уровня, умноженной на $1/(2\sqrt{3})$. Отношение размаха сигнала к средне-квадратичному шуму равно, следовательно,

$$R = 2\sqrt{3}b^n,$$

так как b^n есть число уровней. Выражая это отношение в децибеллах, получаем

$$20 \log_{10} R = 20 \log_{10} 2\sqrt{3} + n(20 \log_{10} b) = 10,8 + n(20 \log_{10} b). \quad (5)$$

¹⁾ Предполагается, что S увеличивается пропорционально W , чтобы скомпенсировать соответствующее увеличение N .

В двоичной системе $b = 2$ и

$$20 \log_{10} R \approx 10,8 + 6n.$$

Рассматривая формулу (5), нужно вспомнить, что n (число цифр) есть множитель, связывающий полную ширину полосы, используемую для передачи, с шириной спектра передаваемого сигнала, т. е. $W = nW_0$. Это подобно индексу модуляции в ЧМ¹⁾. Каждый раз, когда полоса получает приращение W_0 , число n может быть увеличено на единицу, а это увеличивает отношение сигнал/шум на постоянное число децибелл. Другими словами, в КИМ отношение сигнал/шум в децибеллах изменяется линейно с числом цифр в кодовой группе, а следовательно с шириной полосы. Разумеется, при увеличении ширины полосы возрастает мощность шума и необходимо соответственное увеличение мощности сигнала для того, чтобы удержаться на прежнем уровне над порогом.

Двоичная система КИМ, применяющая полосу в десять раз более широкую, чем спектр первоначального сигнала, дает отношение сигнал/шум, равное 70 дб. Системы с более высоким основанием требуют меньшей ширины полосы.

Помехоустойчивость

Важной характеристикой системы передачи является ее восприимчивость к шумам. В системе КИМ шум не оказывает влияния при условии, что пиковое значение не превосходит половины интервала между уровнями импульсов. В двоичной («да — нет») системе это соответствует половине высоты импульса. Аналогично, такие шумы, как посторонние импульсы или влияние импульсов соседних каналов, не окажут влияния, если пиковое значение этих помех плюс пиковое значение шума не превосходят половины высоты импульса. Таким образом, наличие этих помех увеличивает порог, необходимый для удовлетворительного действия системы. Однако если предусмотрено соответствующее превышение над порогом, то наличие сравнительно сильного шума может не оказывать никакого влияния на работу схемы. Поэтому система КИМ, и в особенности двоичная, высокопомехоустойчива.

Когда несколько радиосвязей сходятся в одном пункте или следуют по одним путям между городами, помехоустойчивость канала приобретает особо важное значение. Если восприимчивость каналов к взаимным помехам велика, то требуется большее число отдельных частотных полос и общая полоса, занимаемая передачей, будет велика. Хотя КИМ требует первоначального увеличения ширины полосы для каждого канала, получаемая помехоустойчивость

¹⁾ Частотная модуляция. — *Прим., ред.*

позволяет осуществить много связей, сходящихся или расходящихся из одного пункта и занимающих одну и ту же полосу частот. На одном и том же пути может применяться разделение двух каналов по плоскости поляризации; направленность практически применяемых антенн также позволяет при небольшом отличии направлений приема разделить две связи, работающие на одной частоте. В результате использование диапазона при КИМ исключительно благоприятно, и остальные ее преимущества реализуются при незначительном увеличении общей ширины полосы.

4. Сравнение КИМ и ЧМ

Одно из свойств КИМ состоит в том, что отношение сигнал/шум может быть увеличено за счет расширения полосы. Это преимущество достигается и в других импульсных системах, а также при ЧМ. Так как ЧМ наиболее известная из этих систем, то интересно сравнить КИМ и ЧМ.

Выигрыш от применения широкой полосы

Применяя ЧМ с большим частотным отклонением, получаем, по сравнению с АМ, выигрыш в отношении напряжений сигнал/шум (при той же мощности и при том же шуме на единицу ширины полосы), пропорциональный индексу модуляции, т. е. отношению половины ширины полосы, фактически используемой при передаче, к ширине спектра передаваемого сигнала. Это отношение соответствует n в наших обозначениях. Если мощность шума равномерно распределена по частоте и если требуется сохранить тот же уровень над порогом для ЧМ с различной шириной полосы, то мощность передатчика должна быть пропорциональна ширине полосы (пропорциональна n). Если изменять таким образом мощность, изменяя ширину полосы ЧМ с большим отклонением, то отношение сигнал/шум будет меняться как $n(n^{1/2})$; где множитель $n^{1/2}$ появляется вследствие увеличения напряжения сигнала. Таким образом, отношение сигнал/шум будет равно

$$R = \text{const} \cdot n^{\frac{3}{2}},$$
$$20 \log_{10} R = 30 \log_{10} n + \text{const}. \quad (6)$$

Для двоичной КИМ из равенства (5) получаем такое же взаимозависимое изменение ширины полосы и мощности

$$20 \log_{10} R = 6n + 10,8.$$

Для троичной (основание 3) КИМ

$$20 \log_{10} R = 9,54n + 10,8.$$

Итак, при возрастании ширины полосы (пропорциональной n) отношение сигнал/шум возрастает для ЧМ как $\log n$, тогда как для КИМ оно возрастает как n . Таким образом, КИМ дает больший выигрыш при расширении полосы. Более того, детальное исследование показывает, что, по крайней мере в идеальном случае, КИМ дает почти такое же отношение сигнал/шум, какое вообще достижимо при какой бы то ни было системе модуляции.

Почему же КИМ оказывается такой выгодной с точки зрения увеличения отношения сигнал/шум за счет расширения полосы? Очень ясное представление об этом можно получить, рассматривая простую систему КИМ, в которой четыре двоичных цифры передаются по четырем прилегающим друг к другу частотным полосам с мощностями, как раз достаточными, чтобы перекрыть шум. На рис. 2, *a* сигналы в этих четырех каналах B_1, B_2, B_3 , и B_4 показаны в зависимости от времени. Заштрихованный прямоугольник изображает импульс, незаштрихованный — отсутствие импульса. Прямоугольники имеют длину $\tau_0 = 1/(2 W_0)$. Последовательность кодовых групп, изображенная на рисунке, представляет собой квантованное приближение к линейному возрастанию напряжения со временем, показанное на рис. 2, *b*.

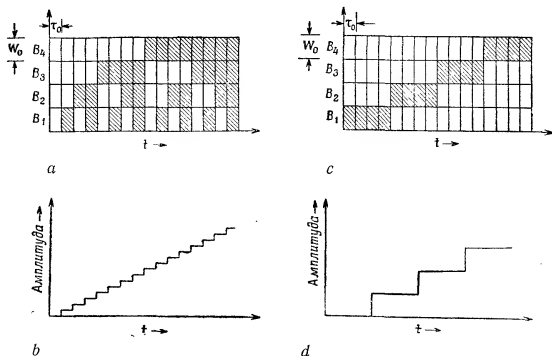
Положим теперь, что передача ограничена посылкой импульса одновременно только по одному каналу, как показано на рис. 2, *c*. Наилучшее квантованное представление сигнала показано на рис. 2, *d*. Здесь число уровней равно четырем, тогда как на рис. 2, *b* имеется 16 уровней. Другими словами, рис. 2, *b* соответствует в четыре раза большему отношению сигнал/шум, чем рис. 2, *d*.

Полная энергия, передаваемая в каждом случае, представляется суммарной заштрихованной площадью; при этом на рис. 2, *a* используется в среднем вдвое большая мощность, чем на рис. 2, *c*. Поэтому получается увеличение отношения сигнал/шум на 12 дБ при увеличении мощности всего на 3 дБ при передаче, соответствующей рис. 2, *a*, по сравнению с передачей, соответствующей рис. 2, *c*. Если рассмотреть случай шести каналов вместо четырех, то получится увеличение отношения сигнал/шум на 21 дБ при увеличении мощности в среднем на 4,77 дБ. Чем больше число каналов и, следовательно, чем шире полоса, тем выгоднее оказывается способ передачи, изображенный на рис. 2, *a*, по сравнению со способом рис. 2, *c*.

Теперь заметим, что рис. 2, *a* изображает КИМ, тогда как рис. 2, *c* изображает по существу квантованную передачу отсчетов при помощи ЧМ. Сигнал на рис. 2, *c* изменяется по частоте в соответствии с амплитудой сигнала. Итак, мы сравнили КИМ и некоторую разновидность ЧМ к очевидному преимуществу КИМ.

Затруднение с ЧМ сигналом рис. 2, *c* состоит в том, что применяется лишь небольшое число возможных сигналов из числа тех, которые могли бы быть посланы по четырем полосам $B_1 \div B_4$;

все остальные, а именно те, для которых сигнал имеется одновременно больше чем в одной полосе, отбрасываются. В идеале КИМ имеет преимущество за счет применения всех возможных сигналов,



Р и с. 2. Сигналы в каналах B_1, B_2, B_3 и B_4 .

(а) Сигнал в системе кодово-импульсной модуляции с частотным разделением. (б) Амплитуды, соответствующие (а). (с) Сигнал в системе квантованной частотной модуляции. (д) Амплитуды, соответствующие (с).

которые могут быть переданы в пределах данной полосы при помощи импульсов с дискретными амплитудами¹).

Соотношение между ЧМ и КИМ очень сходно с соотношением между двумя типами вычислительных машин: так называемыми аналоговыми машинами и цифровыми машинами. В аналоговых машинах числа пропорциональны некоторой физической величине, способной изменяться непрерывно. Типичными примерами являются счетная линейка, электроинтегратор. Увеличение точности требует, вообще говоря, пропорционального увеличения пределов изменения физических величин, используемых для представления

¹) Можно заметить, что возможны сигналы с более тонкой структурой по частоте, чем показанные на рис. 2, а. Это возможно, только если τ сделано больше, так что импульсы, представляющие отсчеты, следуют реже, имеют большую ширину и более узкий спектр. Это означает уменьшение W_0 .

чисел. Более того, малые ошибки накапливаются и не могут быть исключены. В цифровых машинах числа выражаются в цифровой форме, а цифры представляются состояниями некоторых физических деталей машины, которые могут принимать одно из конечного числа возможных состояний. Типичными цифровыми машинами являются счеты, арифмометр и электронные вычислительные машины. В этом типе машин точность растет экспоненциально с числом цифр, а следовательно, и с размерами машины. Малые ошибки, недостаточные для перевода какой-либо детали из одного состояния в другое, не оказывают влияния и не накапливаются.

Аналогично этому в ЧМ амплитуда звукового сигнала измеряется радиочастотой. Для увеличения точности вдвое требуется, грубо говоря, увеличение полосы качания, а следовательно, и ширины полосы также вдвое. В КИМ удвоение ширины полосы позволяет удвоить число цифр и таким образом не удвоить, а возвести в квадрат число различных уровней.

Прочие факторы

Имеются, однако, и другие соображения по поводу сравнения КИМ и обычной некантованной ЧМ. Так, например, КИМ требует применения восстанавливающих повторителей, а ЧМ не требует. КИМ естественно применима, как и другие импульсные системы, к многоканальным системам с временным разделением. С другой стороны, когда при хорошем приеме принятый сигнал значительно превышает порог, отношение сигнал/шум улучшается при ЧМ, но не при КИМ. Если рассматривать передатчики и приемники, то оказывается, что сейчас по крайней мере для больших отношений сигнал/шум аппаратура ЧМ несколько проще, чем аппаратура КИМ.

5. Заключение

КИМ дает большее увеличение отношения сигнал/шум, чем другие системы, вроде ЧМ, также использующие широкую полосу.

При применении двоичной («да — нет») КИМ сигнал высокого качества может быть получен в таких плохих условиях в смысле шума и помех, что едва возможно распознать наличие каждого импульса. Более того, при применении восстанавливающих повторителей, обнаруживающих наличие или отсутствие импульса, а затем воссоздающих импульсы исправленной формы и с надлежащим расположением во времени, первоначальное отношение сигнал/шум может быть сохранено на протяжении длинной цепи повторителей.

КИМ естественно ведет к многоканальной системе с временным разделением.

КИМ не дает увеличения отношения сигнал/шум в течение периодов, когда сигнал силен или шум мал.

КИМ передатчики и приемники несколько сложнее, чем применяемые при других видах модуляции.

В общем представляется, что КИМ идеальным образом подходит для многоканальных систем связи, в которых требуется стандартное качество и высокая надежность.

Приложение I

Покажем, что функция времени $f(t)$, не содержащая составляющих с частотами выше W_0 герц, вполне определяется значениями $f(t)$ в точках отсчета, расположенными на $1/(2W_0)$ секунд друг от друга. Пусть комплексный спектр функции будет

$$F(\omega) = \int_{-\infty}^{+\infty} e^{-i\omega t} f(t) dt.$$

По условию

$$F(\omega) = 0 \quad \text{при } |\omega| > 2\pi W_0.$$

На интервале $(-2\pi W_0, +2\pi W_0)$ можно разложить $F(\omega)$ в ряд Фурье с коэффициентами

$$a_n = \frac{1}{4\pi W_0} \int_{-2\pi W_0}^{2\pi W_0} F(\omega) e^{-i\omega n \frac{1}{2W_0}} d\omega. \quad (1)$$

Так как $F(\omega)$ есть преобразование Фурье для $f(t)$, то $f(t)$ есть обратное преобразование Фурье для $F(\omega)$:

$$f(t) = \frac{1}{2\pi} \int_{-\infty}^{+\infty} F(\omega) e^{i\omega t} d\omega = \frac{1}{2\pi} \int_{-2\pi W_0}^{2\pi W_0} F(\omega) e^{i\omega t} d\omega,$$

так как $F(\omega)$ равно нулю вне указанных пределов.

Если положить $t = \frac{n}{2W_0}$, получим

$$f\left(\frac{n}{2W_0}\right) = \frac{1}{2\pi} \int_{-2\pi W_0}^{2\pi W_0} F(\omega) e^{i\omega \frac{n}{2W_0}} d\omega. \quad (2)$$

Сравнивая (1) и (2), видим, что

$$a_n = \frac{1}{2W_0} f\left(-\frac{n}{2W_0}\right).$$

Итак, если функция известна в точках отсчета

$$\dots -\frac{2}{2W_0}, -\frac{1}{2W_0}, 0, \frac{1}{2W_0}, \frac{2}{2W_0} \dots,$$

то коэффициенты a_n определены. Эти коэффициенты определяют спектр $F(\omega)$, а $F(\omega)$ определяют $f(t)$ для всех значений t . Это значит, что существует единственная функция, не содержащая частот выше W_0 и проходящая через заданные мгновенные значения в точках отсчета, отстоящих друг от друга на $1/(2W_0)$.

Для восстановления функции по этим мгновенным значениям заметим, что

$$F(\omega) = \sum_n a_n e^{i \frac{n\omega}{2W_0}} \quad \text{при } |\omega| < 2\pi W_0,$$

$$F(\omega) = 0 \quad \text{при } |\omega| > 2\pi W_0.$$

Взяв обратное преобразование, получаем

$$\begin{aligned} f(t) &= 2W_0 \sum_n a_n \frac{\sin \pi (2W_0 t + n)}{\pi (2W_0 t + n)} = \sum_n f\left(-\frac{n}{2W_0}\right) \frac{\sin \pi (2W_0 t + n)}{\pi (2W_0 t + n)} = \\ &= \sum_n f\left(\frac{n}{2W_0}\right) \frac{\sin \pi (2W_0 t - n)}{\pi (2W_0 t - n)}. \end{aligned}$$

Другими словами, функцию $f(t)$ можно представить в форме суммы элементарных функций вида $\frac{\sin x}{x}$, расположенных симметрично относительно точек отсчета и имеющих пиковое значение, равное $f(t)$ в соответствующих точках отсчета. Для восстановления функции $f(t)$ нужно лишь создать серию импульсов $\frac{\sin x}{x}$, пропорциональных отсчетам, и сложить их.

Приложение II

Найдем среднюю мощность серии импульсов вида

$$f(t) = \frac{\sin \pi \frac{t}{\tau}}{\pi \frac{t}{\tau}},$$

следующих с неизменной частотой $\frac{1}{\tau}$.

Сигнал может быть записан в виде

$$v(t) = \sum_{k=1}^n V_k f(t - k\tau),$$

где V_k — пиковая амплитуда импульса, происходящего в момент

$t = k\tau$. Средняя «мощность» (т. е. средне-квадратичное значение) S сигнала будет

$$S = \bar{v}^2 = \lim_{n \rightarrow \infty} \frac{1}{n\tau} \int_{-\infty}^{+\infty} v^2(t) dt = \lim_{n \rightarrow \infty} \frac{1}{n\tau} \left| \sum_{k=1}^n V_k^2 \int_{-\infty}^{+\infty} f^2(t - k\tau) dt + \right. \\ \left. + \sum_{j=1}^n \sum_{k=1}^n V_j V_k \int_{-\infty}^{+\infty} f(t - k\tau) f(t - j\tau) dt \right|_{j \neq k}.$$

Для взятой нами записи импульса первый интеграл равен τ , тогда как второй равен нулю. Итак,

$$S = \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{k=1}^n V_k^2.$$

S — это просто средний квадрат отдельных пиковых значений импульсов и может быть записан в виде

$$S = \int_{-\infty}^{+\infty} V^2 p(V) dV,$$

где $p(V) dV$ — вероятность того, что амплитуда импульса лежит в интервале между V и $V + dV$.

Предположим, что импульсы имеют b дискретных уровней амплитуд, различающихся на $K\sigma$ и меняющихся от 0 до $(b-1)K\sigma$. Тогда каждый импульс имеет высоту $aK\sigma$, где a — целое число. Средняя мощность будет равна

$$S = K^2 \sigma^2 \sum_{a=0}^{a=b-1} a^2 p(a),$$

где $p(a)$ — вероятность уровня a . Если все уровни равновероятны, то $p(a) = \frac{1}{b}$ и

$$S = K^2 \sigma^2 \frac{1}{b} \sum_{a=0}^{a=b-1} a^2 = K^2 \sigma^2 \frac{1}{6} (b-1)(2b-1).$$

Величина

$$\frac{1}{b} \sum_{a=0}^{b-1} a^2$$

есть квадрат радиуса инерции (т. е. средне-квадратичный радиус) относительно конца для совокупности b точек, линейно расположенных на единичном расстоянии друг от друга. Средняя мощность любого распределения амплитуд равна среднему из квад-

ратов высот и, следовательно, пропорциональна квадрату радиуса инерции. Радиус инерции относительно некоторой точки есть

$$r^2 = r_0^2 + d^2,$$

где r — радиус инерции относительно выбранной точки, r_0 — радиус инерции относительно центра тяжести, d — расстояние от выбранной точки до центра тяжести. Очевидно, $r_0 < r$, так что средняя мощность будет наименьшей, когда средняя высота равна нулю. S будет наименьшей, если высоты изменяются в пределах

$$\text{от } -K\sigma \frac{1}{2}(b-1) \quad \text{до} \quad +K\sigma \frac{1}{2}(b-1)$$

и выражается в этом случае как

$$S = K^2 \sigma^2 \left| \frac{(b-1)(2b-1)}{6} - \left(\frac{b-1}{2} \right)^2 \right| = K^2 \sigma^2 \frac{b^2-1}{12}.$$

Это можно записать также в виде $S = A^2(b+1)/12(b-1)$, где A — полный диапазон амплитуд, равный $(b-1)K\sigma$; при $b \rightarrow \infty$ $S \rightarrow A^2/12$. Итак, если все амплитуды в диапазоне возможны и равновероятны, средне-квадратичная амплитуда распределения будет равна $\sqrt{S} = A/(2\sqrt{3})$.

ЛИТЕРАТУРА

1. Goodall W. M., Telephony by pulse code modulation, *BSTJ*, **26**, July (1947), 395—409.
2. Grieg D. D., Pulse count modulation system, *Tele-Tech.*, **6**, Sept, (1947), 48—52.
3. Grieg D. D., Pulse count modulation, *El. Commun.*, **24**, Sept. (1947), 287—296.
4. Black H. S., Edson J. O. PCM equipment. *El. Eng.* **66**, Nov. (1947), 1123—1125.
5. Clavier A. C., Grieg D. D., Panter P. F., PCM distortion analysis. *El. Eng.*, **66**, Nov. (1947), 1110—1120.
6. Meacham L. A., Peterson E., An experimental multichannel pulse code modulation system of toll quality, *BSTJ*, **27**, Jan. (1948), 1—43.

Содержание

Разработан метод геометрического представления системы связи. Сообщения и соответствующие им сигналы изображаются точками в двух функциональных пространствах: модуляция есть отображение одного пространства в другое. На основе этих представлений выводится ряд результатов общей теории связи, относящихся к сжатию и расширению полосы частот и к пороговому эффекту. Получены формулы для максимальной скорости передачи двоичных символов по системе, в которой сигнал подвержен действию различного рода помех. Обсуждены некоторые свойства «идеальной» системы, в которой передача ведется с максимальной скоростью. Подсчитано эквивалентное число двоичных символов в секунду для некоторых источников информации.

1. Введение

Общая система связи изображена схематически на рис. 1. Она состоит в основном из пяти частей.

1. Источник информации. Источник избирает одно из совокупности возможных сообщений, подлежащее передаче на приемный конец.

Сообщения могут быть различных типов: например, последовательность букв или цифр, как в телеграфии, или непрерывная функция времени $f(t)$, как при телефонии или радиовещании.

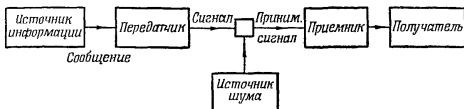
2. Передатчик, который обрабатывает определенным образом сообщения и производит сигнал, могущий быть переданным по каналу на приемный конец.

В телефонии эта операция сводится просто к преобразованию звукового давления в пропорциональный ему электрический ток. В телеграфии имеется операция кодирования, дающая последовательность точек, тире и пауз, соответствующая буквам сообщения. В более сложном примере многоканальной КИМ - телефонии

¹⁾ Shannon C., Communication in the presence of noise, *PIRE*, 37, 1 (1949), 10.

для образования сигнала различные речевые сообщения должны быть отсчитаны, сжаты, квантованы и кодированы и, наконец, надлежащим образом взаимно расположены.

3. Канал — это просто среда, используемая для передачи сигнала от передающего конца к приемному. Это может быть пара проводов, коаксиальный кабель, полоса радиочастот и т. д. Во время передачи или на приемном конце сигнал может искажаться и на него может налагаться помеха. Искажения и влияния помехи могут быть разделены на том основании, что искажения есть результат



Р и с. 1. Общая система связи.

определенной операции, производимой над сигналом, тогда как помеха производит статистические изменения сигнала, которые нельзя предвидеть. Искажения в принципе можно скорректировать обратной операцией; изменения же сигнала, обусловленные помехой, не всегда могут быть устранены, так как сигнал не всегда подвергается при передаче одинаковым изменениям.

4. Приемник. Приемник преобразовывает принятый сигнал и предназначен для восстановления по нему первоначального сообщения. Обычно действие приемника в математическом смысле обратное действию передатчика, хотя могут быть и некоторые различия, когда при конструировании стремятся подавить помеху.

5. Получатель — лицо или устройство, для которого предназначено сообщение.

Следуя Найквисту¹⁾ и Хартли²⁾, принимают логарифмическую меру информации. Если устройство имеет n возможных состояний, оно может по определению хранить $\log_b n$ единиц информации. Выбор основания b ведет к выбору единицы, так как $\log_b n = \log_b c \log n$. Здесь выбрано основание 2 и полученные единицы названы двоичными единицами, или битами. Группа из m реле или триггеров имеет 2^m возможных комбинаций состояний и может поэтому хранить $\log_2 2^m = m$ двоичных единиц.

¹⁾ Nyquist H., Certain factors affecting telegraph speed, *BSTJ*, 3, Apr. (1924), 324.

²⁾ Hartley R., The transmission of information, *BSTJ*, 7, July (1928), 535. (Русский перевод Хартли Р., Передача информации, в сб. Теория информации и ее приложения, Гостехиздат, М., 1959, 5. — Прим. ред.)

Если возможно надежно различить в канале M функций сигнала длительностью T , то можно сказать, что канал может передавать $\log_2 M$ битов за время T . Скорость передачи будет $(1/T) \log_2 M$. Точнее, пропускная способность канала может быть определена как

$$C = \lim_{T \rightarrow \infty} \frac{\log_2 M}{T}. \quad (1)$$

Требованию надежной различимости будет придан позднее точный смысл.

2. Теорема отсчетов¹⁾

Пусть канал пропускает полосу частот W (в герцах, считая от нуля) и может использоваться на протяжении времени T . Без дальнейших ограничений это значит, что можно применять в качестве функций сигнала любые функции, спектр которых полностью помещается в полосу W и длительность укладывается в интервал T . Хотя оба условия не могут быть выполнены в точности, возможно ограничить спектр полосой W и при этом иметь очень малые значения функции вне интервала T . Можно ли более удобным образом описать функции, удовлетворяющие этим условиям? Ответ дает следующая теорема:

Теорема 1. Если функция не содержит частот выше W гц, она полностью определяется своими мгновенными значениями в моменты, отстоящие друг от друга на $1/(2W)$ сек.

Это общеизвестный в теории связи факт. Интуитивное подтверждение состоит в том, что если $f(t)$ не содержит частот выше W , то она не может существенно изменить свое значение за время меньшее, чем половина периода наивысшей частоты, т. е. $1/(2W)$. Математическое доказательство, показывающее, что это положение верно не только приблизительно, но в точности, состоит в следующем.

Пусть $F(\omega)$ есть спектр $f(t)$. Тогда

$$f(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} F(\omega) e^{i\omega t} d\omega = \quad (2)$$

$$= \frac{1}{2\pi} \int_{-2\pi W}^{2\pi W} F(\omega) e^{i\omega t} d\omega, \quad (3)$$

так как $F(\omega)$ равно нулю вне полосы W . Положив

$$t = \frac{n}{2W}, \quad (4)$$

¹⁾ Эта теорема известна в советской литературе как теорема Котельникова. — Прим. ред.

где n — любое целое число, положительное или отрицательное, получим

$$f\left(\frac{n}{2W}\right) = \frac{1}{2\pi} \int_{-2\pi W}^{2\pi W} F(\omega) e^{i\omega \frac{n}{2W}} d\omega. \quad (5)$$

Слева записаны значения $f(t)$ в моменты отсчетов. В интеграле справа можно узнать n -й коэффициент разложения функции $F(\omega)$ по-гармоникам с периодом от $-W$ до $+W$. Это означает, что отсчеты $f(n/2W)$ определяют коэффициенты Фурье в разложении $F(\omega)$. Значит, они определяют и $F(\omega)$, так как $F(\omega)$ равна нулю при всех частотах выше W , а для частот ниже W , $F(\omega)$ определяется своими коэффициентами Фурье. Но $F(\omega)$ полностью определяет исходную функцию $f(t)$, так как функция определена, если ее спектр известен. Итак, исходные отсчеты полностью определяют функцию $f(t)$. Существует одна и только одна функция с ограниченным полосой W спектром и принимающая заданные значения в моменты отсчетов, отстоящие на $1/(2W)$ друг от друга. Функция может быть просто восстановлена по отсчетам, если применить импульс вида

$$\frac{\sin 2\pi W t}{2\pi W t}. \quad (6)$$

Эта функция равна единице при $t = 0$ и нулю при $t = n/(2W)$, т. е. во всех остальных точках отсчета. Ее спектр равен постоянной в полосе W и нулю вне этой полосы. В каждой точке отсчета помещается такой импульс с амплитудой, равной отсчету в данной точке. Сумма импульсов и дает исходную функцию, так как она удовлетворяет условию ограничения спектра и принимает заданные значения.

Математически процесс описывается следующим образом. Пусть x_n есть n -й отсчет. Тогда функция $f(t)$ представляется как

$$f(t) = \sum_{n=-\infty}^{\infty} x_n \frac{\sin \pi (2Wt - n)}{\pi (2Wt - n)}. \quad (7)$$

Аналогичный результат получается и в том случае, когда полоса W начинается не от нулевой частоты, что может быть доказано линейным переносом (физически соответствующим однополосной модуляции) предыдущего случая. В этом случае элементарный импульс получается из $\sin x/x$ посредством однополосной модуляции.

Если функция ограничена временным интервалом T и отсчеты отстоят на $1/(2W)$, то всего в интервале T будет $2TW$ отсчетов. Все отсчеты вне интервала T равны нулю, точнее, будем говорить по определению, что функция ограничена интервалом T тогда и только тогда, когда отсчеты вне интервала в точности равны нулю. Тогда можно сказать, что любая функция, ограниченная полосой W

и временным интервалом T , может быть полностью определена заданием $2TW$ чисел.

Теорема 1 была первоначально дана в других формах математиками¹⁾, но, несмотря на ее очевидную важность, не приводилась в литературе по теории связи. Впрочем, Найквист²⁾, а в последнее время Габор³⁾ отмечали, что достаточно приблизительно $2TW$ чисел, основывая свою аргументацию на разложении функции в ряд Фурье на интервале T . Это дает TW синусов и $(TW + 1)$ косинусов вплоть до частоты W . Небольшое расхождение обусловлено тем, что получаемые таким путем функции не точно ограничены полосой W ; вследствие внезапного начала и окончания гармонических составляющих появятся некоторые составляющие вне полосы. Найквист отметил важное значение интервала $1/(2W)$ для телеграфии; здесь этот интервал будет называться интервалом Найквиста, соответствующим полосе W .

Числа в количестве $2TW$, определяющие функцию, не обязательно должны представлять собой равноотстоящие отсчеты. Например, отсчеты могут браться через неравные интервалы, хотя при наличии значительных отклонений отсчеты должны быть известны с большой точностью для правильного восстановления функции. Можно показать, что достаточно знать значение функции и ее производной в точках отсчета, взятых через одну. Значение функции и двух ее производных в каждой третьей точке образуют иную систему параметров, определяющих функцию. Вообще говоря, любая совокупность $2TW$ независимых чисел, связанных с функцией, может применяться для ее описания.

3. Геометрическое представление сигналов

Совокупность трех чисел x_1 , x_2 и x_3 независимо от их происхождения всегда может быть представлена как координаты точки в трехмерном пространстве. Аналогично $2TW$ равностоящих отсчетов сигнала можно представить себе в качестве координаты точки в пространстве $2TW$ измерений. Каждый частный выбор этих чисел соответствует определенной точке в этом пространстве. Таким образом, каждому сигналу с полосой W и длительностью T соответствует ровно одна точка.

¹⁾ Whittaker J., Interpolatory function theory. Cambridge tracts in mathematics and math. physics, № 33, Cambr. Univ. Press, IV, 1935.

²⁾ Nyquist H., Certain topics in telegraph transmission theory, *AIEE Transaction*, Apr. (1928), 617; Bennett W., Time division multiplex systems, *BSTJ*, 20, Apr. (1941), 199. (В работе Беннета установлен результат, аналогичный теореме 1, но на основе стационарного режима.—Прим. ред.)

³⁾ Gabor D., Theory of communication, *JIEE*, 93, part 3 (1946), 429.

Число измерений $2TW$, вообще говоря, очень велико. Телевизионный сигнал с полосой 5 Мгц , продолжающийся один час, будет представляться точкой в пространстве $2 \cdot 5 \cdot 10^6 \cdot 60^2 = 3,6 \cdot 10^{10}$ измерений. Не надо говорить, что такое пространство не может быть изображено наглядно. Возможно, однако, изучать свойства n -мерного пространства аналитически. Эти свойства в значительной степени являются простым обобщением свойств двух- и трехмерного пространств и могут быть найдены рассуждением по индукции. Преимущество геометрического представления сигнала состоит в том, что можно применить язык и соотношения геометрии к проблемам связи. По существу здесь заменяется сложный объект (например, телевизионный сигнал) в простом окружении [достаточно плоскости, чтобы представить сигнал как $f(t)$ простым объектом (точка) в сложном окружении (пространство $2TW$ измерений).

Если представить себе, что все $2TW$ координатных осей взаимно перпендикулярны, то расстояния в пространстве получают простую интерпретацию. Расстояние от начала координат до данной точки по аналогии со случаями двух и трех измерений равно

$$d = \sqrt{\sum_{n=1}^{2TW} x_n^2}, \quad (8)$$

где x_n — n -й отсчет. Далее, так как

$$f(t) = \sum_{n=1}^{2TW} x_n \frac{\sin \pi (2TW - n)}{\pi (2TW - n)}, \quad (9)$$

то

$$\int_{-\infty}^{\infty} f^2(t) dt = \frac{1}{2W} \sum x_n^2 \quad (10)$$

на том основании, что

$$\int_{-\infty}^{\infty} \frac{\sin \pi (2TW - m)}{\pi (2TW - m)} \frac{\sin \pi (2TW - n)}{\pi (2TW - n)} dt = \begin{cases} 0, & m \neq n. \\ \frac{1}{2W}, & m = n, \end{cases} \quad (11)$$

Итак, квадрат расстояния до данной точки есть умноженная на $2W$ энергия (точнее, энергия в единичном сопротивлении) соответствующего сигнала

$$d^2 = 2WE = 2WTP, \quad (12)$$

где P — средняя мощность за время T . Аналогично расстояние между двумя точками есть умноженное на $\sqrt{2WT}$ средне-квадратичное расхождение между двумя соответствующими сигналами.

Если рассматривать только сигналы, средняя мощность которых меньше P , то им будут соответствовать точки внутри сферы радиуса

$$r = \sqrt{2WTR}. \quad (13)$$

Если к сигналу при передаче добавляется помеха, то это означает, что точка сигнала получила в пространстве смещение в некотором направлении, пропорциональное средне-квадратичному значению помехи. Помеха создает небольшую область неопределенности около каждой точки пространства. Определенное же искажение в канале соответствует искривлению пространства, так что каждая точка смещается, но вполне определенным образом.

В обычном трехмерном пространстве можно построить много различных координатных систем. Эта возможность сохраняется в рассматриваемом многомерном пространстве сигналов. Различные координатные системы соответствуют различным способам описания той же функции сигнала. Примерами могут служить приведенные выше способы определения функции. Другой способ, важный в вопросах связи, состоит в представлении функции ее частотными составляющими. Функция $f(t)$ может быть представлена суммой синусов и косинусов частот, отстоящих на $1/T$, и коэффициенты этого разложения могут быть использованы в качестве некоторой другой совокупности координат. Можно показать, что эти координаты взаимно перпендикулярны и, что существенно, получаются из исходной координатной системы путем поворота.

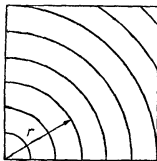
Пропускание сигнала через идеальный фильтр соответствует проектированию точки сигнала на некоторую область в пространстве. Действительно, в частотной системе координат сохраняются координаты, лежащие в полосе пропускания фильтра, остальные же устраняются, так что проектирование происходит на одну из координатных осей, плоскостей или гиперплоскостей. Всякий фильтр осуществляет линейное преобразование векторов пространства, образуя новые векторы, линейно связанные со старыми.

4. Геометрическое представление сообщений

Выше пространство $2TW$ измерений было связано с совокупностью возможных сигналов. Подобным же образом можно связать некоторое пространство с совокупностью возможных сообщений. Предположим, что рассматривается передача речи, так что сообщения состоят из всех возможных звуков, не содержащих частот выше W_1 и длящихся время T_1 .

Точно так же, как сигналы, эти сообщения могут быть однозначным образом изображены в пространстве $2T_1W_1$ измерений. Необходимо, однако, сделать несколько замечаний. Во-первых, различные точки могут изображать одинаковые сообщения, по крайней мере

для получателя. Так, например, в случае речи ухо в известной мере нечувствительно к фазовым искажениям. Сообщения, различающиеся (до некоторой степени) только фазами своих составляющих, звучат одинаково. Это означает возможность сокращения числа существенных измерений пространства сообщений. Все точки, эквивалентные для получателя, могут быть сгруппированы и рассматриваться как одна точка. Тогда может потребоваться меньше чисел для определения каждого из этих «классов эквивалентности», чем для определения каждой отдельной точки. На рис. 2, к примеру,



Р и с. 2. Уменьшение числа измерений посредством эквивалентных групп связи.

изображено двумерное пространство, совокупность точек в квадрате. Если все точки на окружности рассматривать как эквивалентные, то совокупность сводится к одномерной — точка может теперь определяться одним числом, а именно радиусом окружности. В случае звука, если бы ухо было совершенно нечувствительно к фазе, то число измерений могло бы быть уменьшено вдвое. Коэффициенты a_n и b_n при синусах и косинусах для данной частоты не нужно было бы задавать независимо; достаточно было бы задать лишь полную амплитуду $\sqrt{a_n^2 + b_n^2}$ для этой частоты. Уменьшение разрешающей способности слуха по частоте с повышением частоты указывает на дальнейшее уменьшение числа измерений. В вокадере в значительной мере используется подобная эквивалентность для звуков речи, во-первых, путем устранения в большой степени фазовой информации, а во-вторых, путем группирования частот в полосы, в особенности в области высоких частот.

В других видах связи может и не существовать такого рода классов эквивалентности. Получатель ощущает любое изменение сообщения в полном пространстве сообщений, имеющем $2T_1W_1$ измерений. Так обстоит, по-видимому, дело в телевидении.

Второе замечание относится к ограничениям, обусловленным самим источником информации. Пространство $2T_1W_1$ измерений

содержит точку для любой функции $f(t)$, ограниченной полосой W_1 и длительностью T_1 . Класс передаваемых сообщений может быть лишь малым подмножеством всей совокупности этих функций. Например, звуки речи производятся голосовым аппаратом человека. Если исключить передачу всех иных звуков, то действительное количество измерений существенно сократится. Аналогичный результат вытекает из вероятностных соображений. Некоторые сообщения возможны, но настолько маловероятны относительно других, что можно в известном смысле пренебречь ими. В телевизионном изображении, например, весьма правдоподобно, что последовательные кадры различаются лишь незначительно. Вероятность того, что данный элемент изображения сохранит свою яркость в последующих кадрах, очень велика. Математический анализ показывает, что можно значительно сократить число измерений пространства сообщений, если T_1 велико.

Сейчас не будем углубляться в эти вопросы; предположим лишь, что с учетом сказанного получено пространство сообщений с числом измерений D , которое будет, конечно, меньше или равно $2T_1W_1$. Во многих случаях использование вышеуказанных возможностей потребовало бы чрезмерного усложнения аппаратуры. Тогда система строится в предположении, что все функции различны и что в источнике информации нет ограничений. В этом случае пространство сообщений имеет полное количество измерений $2T_1W_1$.

5. Геометрическое представление передатчика и приемника

Рассмотрим теперь работу передатчика с геометрической точки зрения. На входе передатчика имеем сообщение, т. е. одну точку пространства сообщений. На выходе передатчика сигнал — одну точку в пространстве сигналов. Какой бы вид кода и модуляции ни применялся, передатчик должен установить определенное соответствие между двумя точками в двух пространствах. Каждая точка в пространстве сообщений должна соответствовать точке в пространстве сигналов и не должно быть двух сообщений, соответствующих одному и тому же сигналу. В противном случае приемник не мог бы определить, какое сообщение передано. Геометрически такое соответствие называется отображением. Передатчик отображает пространство сообщений в пространство сигналов.

Аналогично приемник отображает пространство сигналов в пространство сообщений. Здесь, однако, возможно отображение нескольких точек в одну точку. Это значит, что несколько различных сигналов в результате демодуляции и декодирования дают одно и то же сообщение. Так, например (при амплитудной модуляции), фаза несущей теряется при демодуляции, и сигналы, различающиеся

только фазой несущей, отображаются в одинаковые сообщения. При частотной модуляции форма сигнала над порогом ограничения не влияет на сообщение. При КИМ возможно значительное искажение принятых импульсов, не оказывающее влияние на работу приемника.

Пока установлено некоторое соответствие между системой связи и геометрическими образами. Это соответствие резюмируется следующей таблицей.

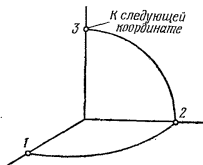
Система связи	Геометрический объект
Совокупность возможных сигналов	Пространство $2TW$ измерений
Отдельный сигнал	Точка в пространстве
Искажение в канале	Искривление пространства
Помеха в канале	Область неопределенности около каждой точки
Средняя мощность сигнала	Умноженный на $(2TW)^{-1}$ квадрат расстояния от начала координат до точки
Совокупность сигналов мощностью P	Совокупность точек в сфере радиуса $\sqrt{2TWP}$
Совокупность возможных сообщений	Пространство $2T_1W_1$ измерений
Совокупность действительных сообщений, различимых получателем	Пространство D измерений, полученное рассмотрением всех эквивалентных сообщений в качестве одной точки и устранением сообщений, которые источник не может дать
Сообщение	Точка в этом пространстве
Передатчик	Отображение пространства сообщений в пространстве сигналов
Приемник	Отображение пространства сигналов в пространстве сообщений

6. Отображение пространства

Можно построить ряд заключений общего характера о способах модуляции, основываясь только на геометрических представлениях. С математической точки зрения простейшими отображениями являются те, в которых оба пространства имеют одинаковые числа измерений. Однополосная амплитудная модуляция является примером такого отображения, и притом особенно простым, так как

координаты в пространстве сигналов пропорциональны соответствующим координатам в пространстве сообщений. При передаче с обоими боковыми полосами пространство сигналов имеет вдвое большее число координат, но они встречаются парами с равными значениями. Если бы пространство сообщений было одномерным, а пространство сигналов двумерным, то это соответствовало бы отображению линии в квадрат, так что точка x на линии отображалась бы точкой (x, x) в квадрате. Таким образом, увеличение числа измерений остается неиспользованным. Все сообщения размещаются в подпространстве, имеющем всего $2T_1W_1$ измерений.

При ЧМ отображение сложнее. Пространство сигналов имеет значительно большее число измерений, чем пространство сообщений. Характер отображения можно пояснить рис. 3, на котором



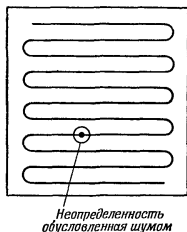
Р и с. 3. Отображение, сходное с ЧМ.

линия отображается в трехмерное пространство. Линия начинается на единичном расстоянии от начала координат на первой оси, переходит, сохраняя это расстояние (т. е. по окружности), на вторую ось, а затем тем же порядком на третью. Легко видеть, что при таком отображении длина линии возрастает пропорционально общему числу координат. Однако ее длина не так велика, как если бы линия, делая петли в пространстве, заполнила весь объем сферы по которой она проходит.

Удлинение линии связано с увеличением отношения сигнал/шум при расширении полосы частот. Так как помеха создает небольшую область неопределенности около каждой точки, то влияние этого обстоятельства на принятое сообщение будет тем меньше, чем больше отображенный образ. Для увеличения образа требуется, чтобы линия переходила взад и вперед в пространстве большего числа измерений, как на рис. 4, на котором изображено отображение линии в квадрат. Нужно заметить, что при этом условии влияние помехи мало по сравнению с длиной линии, если только помеха меньше некоторого критического значения. При этом значении на

приемном конце возникает неуверенность в отношении того, к какой части линии относится сообщение. Это положение имеет общий характер; можно показать, что всякая система, стремящаяся к полному использованию возможностей, связанных с расширением полосы, будет при наличии помех подвержена пороговому эффекту. При малой помехе искажение будет незначительно, но при некотором критическом значении помехи искажение делается очень большим. Это хорошо известно для КИМ.

Предположим теперь, что требуется сократить число измерений, т. е. полосу частот, или время, или и то и другое. Иначе говоря, надо



Р и с. 4. Эффективное преобразование линий в квадрат.

передать сообщения с полосой W_1 и длительностью T_1 по каналу, для которого $TW < T_1W_1$. Было отмечено, что действительное число измерений D пространства сообщений может быть меньше $2T_1W_1$ благодаря свойствам источника и получателя. Следовательно, не требуется, конечно, больше чем D измерений в пространстве сигналов для осуществления хорошего отображения. Чтобы получить эту экономию, нужно выделить эффективные координаты в пространстве сообщений и только их и передавать. Сокращение полосы при передаче речи посредством вокадера может служить примером такой системы.

Возникает вопрос о возможности дальнейшего сокращения. С точки зрения геометрической аналогии — возможно ли отобразить пространства большего числа измерений в пространство меньшего числа измерений? Ответ является, с некоторыми оговорками, положительным. Например, точки квадрата определяются двумя координатами, которые могут быть записаны в десятичной

форме как

$$\left. \begin{aligned} x &= 0 \ a_1 \ a_2 \ a_3 \ \dots \\ y &= 0 \ b_1 \ b_2 \ b_3 \ \dots \end{aligned} \right\} \quad (14)$$

Из этих двух чисел можно построить одно, беря цифры поочередно из x и y

$$z = 0 \ a_1 b_1 a_2 b_2 a_3 b_3 \ \dots \quad (15)$$

По x и y определяем z , и, наоборот, z определяет как x , так и y . Таким образом, здесь имеется однозначное соответствие между точками квадрата и точками на линии.

Этот тип отображений, принадлежащий математике Кантору, может быть легко расширен в направлении сколь угодно большого уменьшения числа измерений. Пространство n измерений может быть взаимно однозначным образом отображено в одномерное. Физически это означает, что произведение полосы и длительности может быть при отсутствии помех сделано сколь угодно малым с сохранением возможности точного восстановления сообщения.

В менее точном смысле отображение, вроде показанного на рис. 4, отображает квадрат в линию, если не требуется точного восстановления начальной точки и мы удовлетворяемся некоторой близлежащей. Чувствительность при увеличении числа измерений, о которой говорилось выше, принимает теперь другой вид. При данном отображении, уменьшающем TW , при изменении сообщения будет некоторый пороговый эффект. Малому изменению сообщения будет соответствовать малое изменение сигнала, пока не будет достигнуто критическое значение. В этой точке сигнал будет претерпевать значительное изменение. В топологии показывается¹⁾, что область пространства большего числа измерений не может быть отображена в область пространства меньшего числа измерений непрерывно. Неизбежный разрыв непрерывности и обуславливает только что описанный нами пороговый эффект в системах связи.

Эти рассуждения связаны с известным «закон Хартли», устанавливающим, что «...верхняя граница могущего быть переданным количества информации определяется суммой произведений из полосы частот на время использования каналов для всех имеющих-ся в распоряжении линий». Это утверждение в одном смысле верно, а в другом нет. Взаимно однозначное и непрерывное отображение (известное в математике как *топологическое* отображение) пространства сообщений в пространство сигналов возможно лишь при условии, что число измерений обоих пространств одинаково, т. е. $D = 2TW$. Следовательно, если ограничиться в передатчике

¹⁾ Hurewicz W., Wallman H., Dimension theory, Princet. Univ. Press., Princeton, 1941. (Русский перевод: Гуревич В., Уолман Х., Теория размерности, ИЛ, М., 1948.—Прим. ред.)

и приемнике однозначными непрерывными отображениями, то произведение TW для канала будет иметь нижнюю границу. Эта нижняя граница определяется не произведением T_1W_1 (длительности и полосы сообщения), а числом существенных измерений D , как указано в разд. 4. Однако нет достаточных оснований к тому, чтобы ограничиваться в передатчике и приемнике топологическими отображениями. Действительно, КИМ и другие подобные системы существенно разрывны и близко соответствуют отображениям вида (14) и (15). Желательно поэтому установить пределы того, что можно сделать без ограничения вида преобразований в приемнике и передатчике. Эти границы, определяемые в последующих разделах, зависят от величины и характера помехи в канале, от мощности передатчика, а также от произведения полосы частот и времени (TW).

Очевидно, что всякая система, увеличивающая или уменьшающая величину TW и полностью использующая дополнительный объем, должна быть существенно нелинейной и довольно сложной в зависимости от особенностей применяемых преобразований.

7. Пропускная способность канала при наличии белого теплового шума

Нетрудно получить некоторые количественные соотношения, получаемые при изменении произведения TW . Предположим, что помеха в системе есть белый тепловой шум, ограниченный полосой W . Эта помеха, добавляясь к переданному сигналу, дает принятый сигнал. Белый тепловой шум характеризуется тем, что изменения каждого отсчета независимы от других и что распределение мгновенных значений подчинено гауссовскому закону со стандартным отклонением $\sigma = \sqrt{N}$, где N — средняя мощность помехи. Сколько различных сигналов можно распознать на приемном конце, несмотря на наличие обусловленных помехой изменений? Грубая оценка может быть получена следующим образом. Если сигнал имеет мощность P , то сигнал, измененный наложенной помехой, будет иметь мощность $P + N$. Число хорошо различимых значений есть

$$K \sqrt{\frac{P+N}{N}}, \quad (16)$$

где K — небольшая константа порядка единицы, зависящая от того, как истолковывается термин «хорошо различимый». Если требуется очень хорошее различие, K будет малым, тогда как, если допускаются случайные ошибки, K будет больше. Так как за время T имеется $2TW$ независимых значений, полное число различимых сигналов будет равно

$$M = \left[K \sqrt{\frac{P+N}{N}} \right]^{2TW}. \quad (17)$$

Число битов, которое можно передать за время T , равно $\log_2 M$ и скорость передачи определяется как

$$\frac{\log_2 M}{T} = W \log_2 K^2 \frac{P + N}{N} \frac{\text{бит}}{\text{сек}}. \quad (18)$$

Трудности, связанные с этим рассуждением, помимо его вообще приблизительного характера, состоят в подразумеваемом предположении, что для различения двух сигналов они должны отличаться в некоторой точке отсчета более чем на ожидаемое значение помехи. Рассуждение предпринимает, что КИМ или нечто весьма сходное с КИМ есть наилучший метод кодирования двоичных цифр и превращения их в сигнал. В действительности два сигнала могут быть надежно различены, если разность между ними мала, но поддерживается неизменной на протяжении значительного промежутка времени. Каждый отсчет принятого сигнала дает при этом малое количество статистической информации о переданном сигнале; в совокупности эти статистические данные дают почти полную достоверность. Эта возможность дает улучшение около 8 дб по сравнению с (18) и позволяет разумно определить уверенную разделимость сигналов, как показано ниже. Теперь воспользуемся геометрическими представлениями для определения точного значения пропускной способности канала при наличии помехи.

Теорема 2. Пусть P — средняя мощность передатчика и пусть помеха есть белый шум с мощностью N в полосе частот W . Применяя достаточно сложную систему кодирования, можно передавать двоичные цифры со скоростью

$$C = W \log_2 \frac{P + N}{N}, \quad (19)$$

со сколь угодно малой частотой ошибок. Никакой метод кодирования не допускает передачи с большей скоростью при произвольно малой частоте ошибок.

Таким образом, величина $W \log (P + N)/N$ вполне определенным образом измеряет способность канала передавать информацию. Это неожиданный результат, так как можно было бы ожидать, что уменьшение частоты ошибок требует соответствующего уменьшения скорости передачи и что последняя должна стремиться к нулю вместе с частотой ошибок. Фактически же можно вести передачу со скоростью C , уменьшая ошибки применением более сложного кодирования и введением более длительных задержек в передатчике и приемнике. На передающем конце будем брать длинную последовательность двоичных цифр и представлять всю такую последовательность в целом специальной функцией сигнала большой длительности. При этом потребуются задержка, так как сигнал определяется только по окончании всей последовательности. Аналогично на

приемном конце передаваемая последовательность может быть восстановлена лишь после того, как принята полностью вся функция соответствующего сигнала.

Докажем теперь теорему 2. В геометрической интерпретации каждая точка сигнала окружена небольшой областью неопределенности, обусловленной помехой. При белом шуме изменения отдельных отсчетов (т. е. координат) являются гауссовскими и независимыми. Таким образом, вероятность смещения с координатами $x_1, x_2, \dots, x_n$ (разностями между координатами принятого и переданного сигналов) равна произведению вероятностей отдельных координат

$$\prod_{n=1}^{2TW} \frac{1}{\sqrt{2\pi 2TW N}} \exp -\frac{x_n^2}{2TW N} = \frac{1}{(2\pi 2TW N)^{TW}} \exp -\frac{1}{2TW} \sum_{n=1}^{2TW} x_n^2.$$

Так как эта величина зависит только от

$$\sum_{n=1}^{2TW} x_n^2,$$

то вероятность данного изменения зависит только от *расстояния*, от исходного сигнала, но не от направления. Другими словами, область неопределенности сферична по своей природе. Хотя для малого числа измерений эта область нерезко ограничена, граница становится все более и более определенной по мере возрастания числа измерений. Это обусловлено тем, что квадрат смещения сигнала равен средней мощности помехи за время T , умноженной на $2TW$. С увеличением T средняя мощность приближается к N . Таким образом, при больших T сигнал будет смещаться в точку, лежащую около сферической поверхности радиуса $\sqrt{2TW N}$ с центром в начале координат. Точнее говоря, если взять T достаточно большим, то с вероятностью, сколь угодно мало отличающейся от 1, смещение будет находиться внутри сферы радиуса $\sqrt{2TW(\bar{N} + \varepsilon)}$, где ε сколь угодно мало. Области неопределенности при очень большом $2TW$ можно себе представить грубо как правильные билиардные шары. Принятые сигналы имеют среднюю мощность $P + N$ и в том же смысле должны почти все лежать на поверхности сферы радиуса $\sqrt{2TW(P + N)}$. Сколько имеется при таких условиях различных сигналов? Ясно, не более чем объем сферы радиуса $\sqrt{2TW(P + N)}$, деленный на объем сферы радиуса $\sqrt{2TW N}$, так как пересечение сфер помехи означало бы смещение сообщений на приемном конце. Объем n -мерной сферы радиуса r равен

$$V = \frac{\pi^{\frac{n}{2}}}{\Gamma\left(\frac{n}{2} + 1\right)} r^n. \quad (20)$$

Таким образом, верхний предел для числа различных сигналов есть

$$M \leq \left(\sqrt{\frac{P+N}{N}} \right)^{2TW}, \quad (21)$$

откуда пропускная способность канала

$$C = \frac{\log_2 M}{T} \leq W \log_2 \frac{P+N}{N}. \quad (22)$$

Это доказывает второе утверждение теоремы.

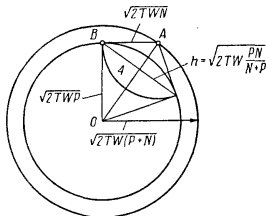
Для доказательства первой части теоремы нужно показать, что существует система кодирования, обеспечивающая передачу со скоростью $W \log_2(P+N)/N$ бит/сек при частоте ошибок, меньшей ϵ , где ϵ произвольно мало. Рассматриваемая система действует следующим образом. Длинная последовательность, содержащая, скажем, m двоичных цифр, воспринимается передатчиком. Всего может быть 2^m таких последовательностей, и каждой из них соответствует специальная функция сигнала длительностью T . Следовательно, имеется $M = 2^m$ различных функций сигнала. Когда данная последовательность из m цифр закончена, передатчик начинает посылать соответствующий сигнал. Приемник принимает искаженный сигнал, сравнивая его с каждым из M возможных переданных сигналов и выбирая тот из них, к которому принятый сигнал ближе всего (в смысле средне-квадратичной ошибки). Затем на выходе приемника воспроизводится соответствующая этому сигналу последовательность двоичных цифр. Таким образом, получается общая задержка на $2T$.

Для обеспечения частоты ошибок, меньшей ϵ , M функций сигнала должны достаточно далеко отстоять друг от друга. Действительно, надо выбрать их так, чтобы, когда принят искаженный сигнал, ближайшая сигнальная точка в геометрической интерпретации соответствовала истинному переданному сигналу с вероятностью, большей чем $1-\epsilon$.

Получается, как ни странно, что достигается наилучший результат, если выбрать M функций сигнала наудачу из числа всех точек внутри сферы радиуса $\sqrt{2TWP}$. Физически это соответствует тому случаю, когда в качестве функций сигнала берутся M различных выборок ограниченного по полосе белого шума.

Некоторый определенный выбор M точек в сфере означает определенный способ кодирования. Общая схема доказательства состоит в рассмотрении всех таких выборов с целью показать, что частота ошибок, усредненная по всем выборам, меньше ϵ . Это покажет, что в совокупности выборов существуют определенные выборы, для которых частота ошибок меньше ϵ . Конечно, будут и другие выборы с большой частотой ошибок.

Геометрические соотношения показаны на рис. 5, представляющем собой плоское сечение многомерной сферы, определяемой типичным переданным сигналом B , принятым сигналом A и центром в O . Переданный сигнал лежит очень близко к поверхности сферы радиуса $\sqrt{2TWP}$, так как в сфере большого числа измерений почти весь объем сосредоточивается около поверхности. Аналогично принятый сигнал лежит на поверхности сферы радиуса $\sqrt{2TW(P+N)}$.



Р и с. 5. Геометрическая интерпретация теоремы 2.

Многомерная чечевицеобразная область L есть область возможных сигналов, породивших A , так как расстояние между переданным и принятым сигналом почти наверняка очень близко к $\sqrt{2TWN}$. Область L имеет меньший объем, чем сфера радиуса h . Можно определить h , приравняв вычисленные двумя различными способами площади треугольника OAB

$$\frac{1}{2} h \sqrt{2TW(P+N)} = \frac{1}{2} \sqrt{2TWP} \sqrt{2TWN},$$

откуда

$$h = \sqrt{2TW \frac{PN}{P+N}}.$$

Вероятность любой определенной сигнальной точки (отличной от истинного сигнала, породившего A), лежащей в L , поэтому меньше, чем отношение объемов сфер с радиусами $\sqrt{2TWP}$ и $\sqrt{2TW(P+N)}$, так как в нашем ансамбле систем кодирования выбраны точки наудачу из числа точек внутри сферы радиуса $\sqrt{2TWP}$. Это отношение равно

$$\left(\frac{\sqrt{2TW \frac{PN}{P+N}}}{\sqrt{2TWP}} \right)^{2TW} = \left(\frac{N}{P+N} \right)^{TW}. \quad (23)$$

Имеется M точек сигнала. Таким образом, вероятность p , что все они, за исключением точки истинного сигнала, находятся вне L , больше чем

$$\left[1 - \left(\frac{N}{P+N} \right)^{TW} \right]^{M-1}. \quad (24)$$

Если точки лежат вне L , сигнал воспроизводится правильно. Поэтому, если сделать p больше чем $1 - \epsilon$, частота ошибок будет меньше ϵ . Это будет справедливо, если

$$\left[1 - \left(\frac{N}{P+N} \right)^{TW} \right]^{M-1} > 1 - \epsilon. \quad (25)$$

Но $(1-x)^n$ всегда больше, чем $1-nx$, если n положительно. Поэтому (25) будет верно, если

$$1 - (M-1) \left(\frac{N}{P+N} \right)^{TW} > 1 - \epsilon \quad (26)$$

или если

$$M-1 < \epsilon \left(\frac{P+N}{N} \right)^{TW}, \quad (27)$$

или же если

$$\frac{\log(M-1)}{T} < W \log \frac{P+N}{N} + \frac{\log \epsilon}{T}. \quad (28)$$

Для любого данного ϵ это неравенство можно удовлетворить, выбирая достаточно большое T . При этом и $\log(M-1)/T$ или $\log M/T$ будет сколь угодно близок к $W \log(P+N)/N$. Это показывает, что при случайном выборе сигнальных точек можно получить сколь угодно малую частоту ошибок и передавать со скоростью, сколь угодно близкой к C . Можно также передавать в точности со скоростью C при произвольно малом ϵ , так как добавочные двоичные цифры вовсе не нужно передавать — они могут случайным образом добавляться на приемном конце. Это лишь добавляет еще одну произвольно малую величину к ϵ . Этим завершается доказательство.

8. Обсуждение

Будем называть систему, обеспечивающую передачу без ошибок со скоростью C , идеальной системой. Такая система не может быть осуществлена ни при каком конечном процессе кодирования, но к ней можно приблизиться настолько, насколько это желательно. По мере приближения к идеалу происходит следующее. 1) Скорость передачи двоичных чисел приближается к $C = W \log_2(1 + P/N)$. 2) Частота ошибок приближается к нулю. 3) Передаваемый сигнал по своим статистическим свойствам приближается к белому шуму. Это справедливо, грубо говоря, потому, что

применяемые функции сигнала должны быть распределены случайно внутри сферы радиуса $\sqrt{2TWP}$. 4) Пороговый эффект становится очень острым. Если помеха превзойдет значение, для которого построена система, частота ошибок возрастает очень быстро. 5) Требуемые задержки в передатчике и приемнике неограниченно возрастают. Конечно, в широкополосной системе задержка в одну миллисекунду может уже рассматриваться как бесконечная.

На рис. 6 построена функция $C/W = \log(1 + P/N)$; по оси абсцисс отложено P/N в децибеллах, по оси ординат отложен C/W ,

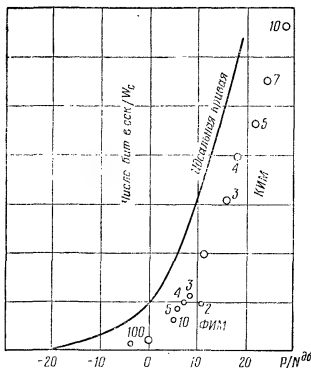


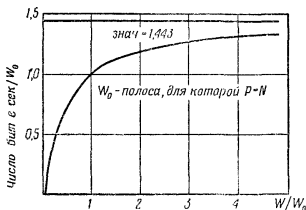
Рис. 6. Сравнение КИМ и ФИМ с идеальной системой.

т. е. число бит на герц. Кружки означают КИМ при двоичном, троичном и т. д. кодах с положительными и отрицательными импульсами, подобранными так, чтобы получалась одна ошибка на 10^5 двоичных цифр. Точки относятся к ФИМ при двух, трех и т. д. дискретных положениях импульса¹⁾. Разница между положением точек и идеальной кривой показывает выигрыш, который мог бы быть

¹⁾ Точки КИМ подсчитаны по формулам, приведенным в статье «Принципы кодово-импульсной модуляции» (см. стр. 414 данного сборника). Точки ФИМ основаны на неопубликованных подсчетах Мак-Миллана, который отмечает, что при очень малых P/N точки приближаются к идеальной кривой с точностью до 3 дБ.

достигнут применением более сложных систем кодирования. Он достигает примерно 8 дБ по мощности для всего практического диапазона. Точки и кружки показывают примерно наилучшие результаты, достижимые без задержки. Стоит ли прибегать к более сложным системам модулирования для частичной реализации возможного выигрыша, — это вопрос экономики.

Величина $TW \log(1 + P/N)$ при большом T выражает число двоичных единиц, которое может быть передано за время T .



Р и с. 7. Пропускная способность канала как функция полосы.

Это выражение можно рассматривать как соотношение обмена между различными параметрами. Величины T , W , P и N могут в отдельности изменяться по желанию без изменения передаваемого количества информации при условии, что величина $TW \log(1 + P/N)$ остается неизменной. Если TW уменьшено, то P/N должно быть увеличено и т. д.

Обычно при увеличении W мощность шума в полосе возрастает пропорционально $N = N_0 W$, где N_0 — мощность шума на один герц. В этом случае имеем

$$C = W \log \left(1 + \frac{P}{N_0 W} \right). \quad (29)$$

Если положить $W_0 = P/N_0$, т. е. определить W_0 как полосу, в которой мощность шума равна мощности сигнала, то (29) примет вид

$$\frac{C}{W_0} = \frac{W}{W_0} \log \left(1 + \frac{W}{W_0} \right). \quad (30)$$

На рис. 7 C/W_0 построено как функция от W/W_0 . При увеличении полосы пропускная способность быстро растет, пока мощность шума примерно не сравняется с мощностью сигнала; после этого пропуск-

ная способность растет все медленнее, приближаясь к асимптотическому значению, равному значению при $W = W_0$, умноженному на $\log_2 e$.

9. Произвольный гауссовский шум

Если белый шум пропустить через фильтр с коэффициентом передачи $Y(f)$, то получится шум со спектром мощности $N(f) = K |Y(f)|^2$, известный под названием гауссовского шума. Можно подсчитать пропускную способность канала, подверженного действию гауссовского шума, основываясь на результатах, полученных для белого шума. Пусть полная мощность передатчика P распределена по спектру по закону $P(f)$. Тогда

$$\int_0^W P(f) df = P. \quad (31)$$

Можно разделить полосу на большое число узких полосок, считая $N(f)$ в каждой полоске постоянным. Полная пропускная способность для данного $P(f)$ будет тогда равна

$$C_1 = \int_0^W \log \left(1 + \frac{P(f)}{N(f)} \right) df, \quad (32)$$

если применить результат для белого шума к каждой полоске. Наибольшая скорость передачи определяется максимизацией C_1 при условии (31). Это значит, что требуется максимизировать

$$\int_0^W \left[\log \left(1 + \frac{P(f)}{N(f)} + \lambda P(f) \right) \right] df. \quad (33)$$

Условие максимума можно получить, используя вариационное исчисление или просто, основываясь на вогнутости кривой $\log(1+x)$; получаем

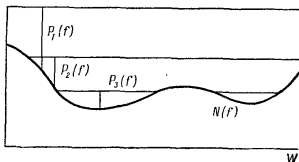
$$\frac{1}{N(f) + P(f)} + \lambda = 0, \quad (34)$$

т. е. $N(f) + P(f)$ должна быть постоянным. Постоянная подбирается так, чтобы полная мощность сигнала равнялась P . На частотах, на которых мощность шума мала, мощность сигнала должна быть велика и наоборот, как и можно было ожидать.

Это проиллюстрировано на рис. 8, где кривая изображает спектр шума и три прямые соответствуют различным P . Если P мало, $P(f) + N(f)$ невозможно сделать постоянным, так как это потребует на некоторых частотах отрицательных мощностей. Легко показать, однако, что в этом случае наилучшее $P(f)$ получается,

если сделать $P(f) + N(f)$ постоянным везде, где это возможно, а на остальных частотах взять $P(f)$ равным нулю. При малых P некоторые частоты окажутся вообще неиспользованными.

Если теперь изменять спектр шума $N(f)$, поддерживая полную мощность шума неизменной и всякий раз подгоняя спектр сигнала $P(f)$ так, чтобы получить наилучшие условия передачи, то можно



Р и с. 8. Наилучшее распределение мощности передатчика.

найти наихудший спектр шума. Это оказывается белый шум. Хотя это показывает лишь, что белый шум — наихудший среди гауссовых шумов, ниже будет показано, что он является наихудшим среди всех возможных шумов с данной мощностью N в полосе.

10. Пропускная способность канала при произвольном типе шума

Существует, конечно, много видов помех, не являющихся гауссовским шумом; таковы, например, импульсная помеха или белый шум, пропущенный через нелинейное устройство. Если на сигнал наложена одна из этих помех, то по-прежнему имеется определенная пропускная способность канала C , максимальная скорость передачи двоичных чисел. Наметим здесь только общую теорию¹⁾. Пусть $x_1, x_2, \dots, x_n$ — мгновенные значения помехи в последовательных точках отсчета и пусть

$$p(x_1, x_2, \dots, x_n) dx_1 \dots dx_n \quad (35)$$

вероятность²⁾ того, что эти значения лежат между x_1 и $x_1 + dx_1$, x_2 и $x_2 + dx_2$ и т. д. Тогда функция p описывает статистическую структуру помехи, поскольку рассматриваются n последовательных отсчетов. Энтропия H помехи определяется следующим

¹⁾ Shannon C., A mathematical theory of communication, *BSTJ*, 27, July, Oct. (1948), 379, 623. (См. стр. 243 данного сборника. — Прим. ред.)

образом. Положим

$$H_n = -\frac{1}{n} \int \dots \int p(x_1, x_2, \dots, x_n) \times \\ \times \log_e p(x_1, x_2, \dots, x_n) dx_1, \dots, dx_n. \quad (36)$$

Тогда

$$H = \lim_{n \rightarrow \infty} H_n. \quad (37)$$

Предел существует во всех практически интересных случаях и во многих из них может быть найден. H есть мера случайности помехи. В случае белого гауссовского шума мощности N энтропия равна

$$H = \log_e \sqrt{2\pi e N}. \quad (38)$$

Удобно измерять случайность произвольного вида помехи не непосредственно ее энтропией, а по сравнению с белым шумом. Можно вычислить мощность белого шума, имеющего такую же энтропию, как данная помеха. Эту мощность, равную

$$\bar{N} = \frac{1}{2\pi e} \exp 2H, \quad (39)$$

где H — энтропия данной помехи, будем называть *энтропийной мощностью* помехи.

Помеха с энтропийной мощностью $\bar{N}$ действует сходно с белым шумом мощностью $\bar{N}$, поскольку рассматривается изменение сообщения. Можно показать, что область неопределенности около каждой точки сигнала будет иметь такой же объем, как при белом шуме. Конечно, это уже не будет сферическая область. При доказательстве теоремы 1 объем области неопределенности был главным использованным признаком помехи. В основном то же рассуждение с небольшими изменениями может быть применено к помехе любого вида. Этот результат выражает следующая теорема.

Т е о р е м а 3. Пусть помеха в ограниченной полосе W имеет мощность N и энтропийную мощность N_1 . Тогда пропускная способность C лежит в пределах

$$W \log_2 \frac{P+N_1}{N_1} \leq C \leq W \log_2 \frac{P+N}{N_1}, \quad (40)$$

где P — средняя мощность сигнала, W — ширина полосы.

Если помеха есть белый шум, то $N_1 = N$, обе границы совпадают, и результат сводится к теореме 2 разд. 7.

Для любой другой помехи $N_1 < N$. Вот почему белый шум есть наихудшая помеха из всех возможных. Если помеха есть гауссов-

ский шум со спектром $N(f)$, то

$$N_1 = W \exp \frac{1}{W} \int_0^W \log N(f) df. \quad (41)$$

Верхняя граница в теореме 3 достигается, когда превзойдена наибольшая мощность помехи на рис. 8. Это легко проверить подстановкой.

В наиболее интересных случаях P/N велико. При этом обе границы примерно равны и можно считать пропускную способность равной $W \log (P + N)/N_1$. Лучше брать верхнюю границу, так как можно показать, что C при увеличении P/N приближается к верхней границе.

11. Дискретные источники информации

До сих пор рассматривались главным образом каналы. Пропускная способность C определяет максимальную скорость передачи случайной последовательности двоичных цифр, если они закодированы наилучшим возможным способом. В общем случае информация, подлежащая передаче, не имеет такой формы. Она может быть, например, последовательностью букв, как в телеграфии, звуками речи или телевизионным сигналом. Можно ли найти эквивалентное число двоичных единиц для источников информации названных типов? Рассмотрим сначала дискретный источник, т. е. случай, когда сообщение состоит из дискретных символов. Вообще говоря, могут существовать разного рода корреляции между отдельными символами. Если сообщение представляет собой английский текст, то буква E встречается чаще всего, за T часто следует H и т. п. Эти корреляции позволяют сжать текст путем соответствующего кодирования. Можно определить энтропию дискретного источника аналогично тому, как это было сделано для шума, а именно

$$H_n = -\frac{1}{n} \sum_{i, j, \dots, s} p(i, j, \dots, s) \log_2 p(i, j, \dots, s), \quad (42)$$

где $p(i, j, \dots, s)$ — вероятность последовательности символов $i, j, \dots, s$, а сумма берется по всем последовательностям из n символов. Тогда энтропия есть

$$H = \lim_{n \rightarrow \infty} H_n. \quad (43)$$

Оказывается, что H есть число двоичных единиц, производимое источником на каждый символ сообщения. Нижеследующая теорема доказывается в приложении.

Теорема 4. *Возможно закодировать все последовательности из n символов сообщения в виде последовательностей двоичных цифр таким образом, что среднее число двоичных единиц на символ сообщения равно приблизительно H , причем приближенное равенство переходит в точное с возрастанием n .*

Отсюда следует, что если имеется канал с пропускной способностью C и дискретный источник с энтропией H , то возможно закодировать сообщения, превратив их через двоичные цифры в сигнал, и вести передачу со скоростью C/H первоначальных символов сообщения в секунду.

Например, если источник производит последовательность букв A, B и C с вероятностью $p_A = 0,6$, $p_B = 0,3$, $p_C = 0,1$ и последовательные буквы выбираются независимо, то $H_n = H_1 = -(0,6 \log_2 0,6 + 0,3 \log_2 0,3 + 0,1 \log_2 0,1) = 1,294$ и информация составляет 1,294 бит на каждую букву сообщения. Канал с пропускной способностью 100 бит/сек. может передавать при наилучшем коде $100/1,294 = 77,3$ букв/сек.

12. Непрерывные источники информации

Если источник производит непрерывную функцию времени, то, если нет оговорок, ему можно приписать бесконечную скорость создания информации. Действительно, для того чтобы охарактеризовать точно величину с непрерывным изменением, требуется бесконечное число двоичных цифр. Невозможно передавать точно непрерывную информацию по каналу с конечной пропускной способностью.

К счастью, нам не нужно передавать непрерывные сообщения в точности. Некоторое расхождение между исходным и восстановленным сообщениями всегда допускается. Если принят известный допуск, то непрерывному источнику может быть приписана определенная конечная скорость в битах/сек. Следует помнить, что эта скорость зависит от природы и величины допускаемой погрешности. Скорость может быть описана как скорость создания информации по отношению к критерию точности.

Положим, что критерий точности есть средне-квадратичное расхождение между исходным и восстановленным сигналами и что можно допустить его значение равным $\sqrt{N_1}$. Тогда каждая точка в пространстве сообщений окружена малой сферой радиуса $\sqrt{2T_1 W_1 N_1}$. Если система такова, что восстановленное сообщение лежит в пределах этой сферы, то передача удовлетворительна. Таким образом, число различимых сообщений, которые должны допускать отдельную передачу, будет порядка отношения объема V_1 области возможных сообщений к объему малых сфер. Проведа

это рассуждение детально аналогично тому, как это делалось в разделах 6 и 9, приходим к следующему.

Теорема 5. Если источник сообщения имеет мощность Q , энтропийную мощность $\bar{Q}$ и ширину полосы W_1 , то скорость R создания информации (в битах/сек) лежит в пределах

$$W_1 \log_2 \frac{\bar{Q}}{N_1} \leq R \leq W_1 \log_2 \frac{Q}{N_1}, \quad (44)$$

где N_1 есть наибольший допустимый средний квадрат ошибки воспроизведения. Если имеется канал с пропускной способностью и источник, создающий информацию со скоростью R , меньшей или равной C , то возможно закодировать источник так, чтобы вести по данному каналу передачу с точностью, определяемой N_1 . Если $R > C$, то это невозможно.

В случае когда источник сообщения производит белый тепловой шум, $\bar{Q} = Q$. При этом оба предела равны $R = W_1 \log Q/N_1$. Возможно, следовательно, передавать белый тепловой шум мощностью Q и с полосой W_1 по каналу с полосой W при наличии помехи в виде белого шума мощностью N и восстановить исходное сообщение с ошибкой, средний квадрат которой равен N_1 , тогда и только тогда, когда

$$W_1 \log \frac{Q}{N_1} \leq W \log \frac{P+N}{N}. \quad (45)$$

Приложение

Рассмотрим возможные последовательности из n символов. Расположим их в порядке убывания вероятности $p_1 \geq p_2 \geq p_3 \dots \geq p_n$. Пусть

$$P_i = \sum_{j=1}^{i-1} p_j.$$

Тогда i -е сообщение кодируется разложением P_i как двоичной дроби с сохранением только первых t_i членов, где t_i определяется соотношением

$$\log_2 \frac{1}{p_i} \leq t_i \leq 1 + \log_2 \frac{1}{p_i}. \quad (46)$$

Наиболее вероятные последовательности имеют короткие коды, маловероятные — длинные коды. Имеем

$$\frac{1}{2^{t_i}} \leq p_i \leq \frac{1}{2^{t_i-1}}. \quad (47)$$

Коды для различных последовательностей будут все различны. P_{i+1} отличается, например, от P_i на p_i , и поэтому его двоичное разложение будет отличаться не менее чем одним из числа первых t_i членов; аналогичное утверждение справедливо для всех других P_i . Средняя длина кодированного сообщения будет $\sum p_i t_i$. Пользуясь неравенством (46), найдем

$$-\sum p_i \log p_i \leq \sum p_i t_i < \sum p_i (1 - \log p_i), \quad (48)$$

или

$$nH_n \leq \sum p_i t_i < 1 + nH_n. \quad (49)$$

Среднее число двоичных цифр на символ сообщения есть $1/n \sum p_i t_i$
и

$$H_n \leq \frac{1}{n} \sum p_i t_i < \frac{1}{n} + H_n.$$

При $n \rightarrow \infty$ $H_n \rightarrow H$ и $1/n \rightarrow 0$, так что среднее число двоичных единиц приближается к H .

Предыдущая работа по теории связи²⁾ показала, что количество информации имеет в рамках теории связи естественную количественную меру, задаваемую формулами типа формулы для энтропии $H = -\sum p \log p$. Это привело к теоремам, дающим наиболее эффективные методы кодирования сообщений, создаваемых стохастическим процессом, в стандартной форме, скажем, в случайную последовательность двоичных знаков, предназначенных для наиболее эффективного использования существующих каналов связи. Однако не было определено понятие информации как таковой. Оказывается возможным сформулировать подход к теории, в которой источники информации в системе, передающей сообщения, являются элементами структуры.

Ведущая идея состоит в том, что любое обратимое преобразование сообщений, создаваемых стохастическим процессом, скажем, посредством невырожденного преобразователя с конечным числом состояний, следует рассматривать как содержащее ту же информацию, что и первоначальное сообщение. С точки зрения теории сообщений знание зашифрованного кодом Морзе текста телеграммы эквивалентно знанию самого текста. Таким образом будем считать информацию источника эквивалентным классом всех обратимых преобразований сообщений, создаваемых источником. Каждое частное преобразование является представителем этого класса, аналогично тому как тензор задается своими компонентами в некоторой координатной системе частного вида.

В зависимости от выбора множества преобразований, рассматриваемых как эквивалентные, могут быть получены различные теории. Два способа выбора ведут к интересным и приложимым на практике результатам: 1) группа всех преобразователей с конечным

¹⁾ Shannon C., Some topics in information theory, Proceedings of the International Congress of Mathematicians, II (1950), 262—263.

²⁾ Shannon C., A mathematical theory of communication. (Русский перевод см. стр. 243 данного сборника.—Прим. пер.).

множеством состояний (эффективно допускающих положительные или отрицательные задержки); 2) группа преобразователей без задержек с конечным множеством состояний, причем требуется, чтобы имеющийся в настоящий момент символ на выходе был функцией от входа в настоящий и прошлый момент и аналогично для обратного преобразователя.

Первый случай — самый простой и связан более тесно с предыдущей работой, в которой допускались неограниченные задержки кодирования в передатчике и приемнике. Транзитивное соотношение включения между элементами информации: $x \geq y$ (указывающее на частичную упорядоченность) означает, что y может быть получен из x с помощью преобразователя с конечным числом состояний (не обязательно обратимого). Энтропия источника (которая инвариантна относительно группы обратимых преобразований) оказывается нормой, монотонно зависящей от порядка. Наименьшая верхняя грань двух элементов является полной информацией об обоих источниках, представлением которой может служить последовательность упорядоченных пар сообщений, получаемых из двух источников. Наибольшая нижняя грань также может быть определена; таким образом, в результате возникает информационная структура. Точная нижняя грань будет существовать всегда, а если множество рассматриваемых источников конечно, то будет существовать и точная верхняя грань. Полученная таким образом структура, вообще говоря, немодулярна. В действительности может быть построена информационная структура, изоморфная любой конечной дискретной структуре.

С помощью формулы $q(x, y) = H_x(y) + H_y(x)$ может быть определена метрика, удовлетворяющая обычным требованиям. Эта метрика вводит топологию — понятие сходящихся в смысле Коши последовательностей информационных элементов и понятие предельной точки. Если сходящиеся последовательности добавлены к структуре как новые точки (с соответствующими модификациями определения равенства и т. д.), то образуются непрерывные структуры, например совокупность всех абстракций¹⁾ полной информации в системе преобразователей с конечными множествами состояний или в предельных последовательностях таких преобразователей.

Теория преобразователей без задержек также ведет к некоторой структуре, но соответствующие задачи, которые, возможно, более важны для приложений, изучены хуже. Энтропии источника теперь уже недостаточно для того, чтобы охарактеризовать источник для целей кодирования, и действительно может быть найдено бесконечное количество независимых инвариантов источника. Некоторые из

¹⁾ Здесь автор использует термин абстракция (abstraction), не разъясняя его смысла в данном контексте. — *Прим. ред.*

них связаны с задачей наилучшего предсказания следующего сигнала, который должен быть произведен при условии, что уже известна вся предыстория. Теория преобразователей без задержек имеет приложение к проблеме передачи сообщения по каналу в случае, когда в нашем распоряжении имеется второй канал для передачи информации в обратном направлении. Второй канал может быть использован в некоторых случаях для улучшения прямой передачи. Для этого случая найдены верхние границы для пропускной способности прямого канала. Теория преобразователей без задержек имеет также приложение к проблеме сглаживания и предсказания по методу наименьших квадратов¹⁾. Фильтр с минимальной фазой имеет обратный фильтр (без задержки) и, следовательно, принадлежит к группе преобразователей без задержек непрерывных временных рядов. Проблема предсказания по методу наименьших квадратов может быть решена при помощи преобразования рассматриваемых временных рядов в каноническую форму и нахождения оператора наилучшего прогноза для этой формы.

¹⁾ B o d e H. W., S h a n n o n C., A simplified derivation of linear least square smoothing and prediction theory. (Русский перевод см. стр. 687 данного сборника.—Прим. пер.)

ПРОПУСКНАЯ СПОСОБНОСТЬ КАНАЛА С ШУМОМ ПРИ НУЛЕВОЙ ОШИБКЕ¹⁾

Краткое содержание

Пропускная способность C_0 канала с шумом при нулевой ошибке определяется как наименьшая верхняя грань скоростей, с которыми возможно вести передачу информации при нулевой вероятности ошибки. Изучаются различные свойства C_0 ; даются оценки сверху и снизу и приводятся методы вычисления C_0 . Получены неравенства для C_0 , связывающие «сумму» и «произведение» двух данных каналов. Рассматриваются аналогичные задачи для пропускной способности канала при нулевой ошибке и при наличии линии обратной связи. Показано, что, тогда как обычная пропускная способность дискретного канала без памяти при наличии обратной связи равна пропускной способности того же самого канала в отсутствии обратной связи, пропускная способность C_{0F} канала с обратной связью при нулевой ошибке может быть больше. Приводится решение задачи вычисления C_{0F} .

Введение

Обычную пропускную способность канала с шумом можно понимать следующим образом. Для данного канала существует некоторая последовательность кодов с увеличивающейся длиной блока, такая, что входная скорость передачи приближается к C , а вероятность ошибки декодирования на приемном конце стремится к нулю. Больше того, это не так для любой величины, превышающей C . В некоторых случаях более интересно рассмотреть не коды с вероятностью ошибки, *приближающейся* к нулю, а коды, для которых вероятность *равна* нулю, и исследовать наибольшую скорость передачи (или наименьшую верхнюю границу этих скоростей), возможную для таких кодов. Эта скорость C_0 составляет главный предмет исследования настоящей статьи. Интересно, что, хотя C_0 описывает казалось бы более простое, по сравнению с C , свойство канала,

¹⁾ Shannon C., The zero error capacity of a noisy channel, *IRE Trans.*, № 3 (1956), 8.

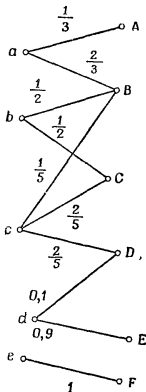
ее в действительности более трудно вычислять и это вычисление связано с рядом еще не решенных проблем.

Здесь будут рассмотрены только конечные дискретные каналы без памяти. Такие каналы определяются конечной матрицей переходов $\|p_i(j)\|$, где $p_i(j)$ — вероятность того, что входная буква i будет принята как выходная буква j ($i = 1, 2, \dots, a$; $j = 1, 2, \dots, b$) и $\sum_j p_i(j) = 1$. В эквивалентной форме такой канал может быть представлен схематически, как это показано на рис. 1.

То, что рассматриваемый канал является каналом *без памяти*, означает, что последовательные операции независимы. Если были использованы входные буквы i и j , то вероятность выходных букв k и l будет $p_i(k)p_j(l)$. Последовательность входных букв будет называться *входным словом*, последовательность выходных букв — *выходным словом*. Отображение M сообщений (которые можно перенумеровать цифрами $1, 2, \dots, M$) в некоторое подмножество входных слов длины n будет называться *блоковым кодом* длины n . *Входной скоростью* этого кода будет называться $R = \frac{1}{n} \log M$. Если не будет специальных оговорок, то слово «код» будет означать такой блочковый код. В работе будут употребляться натуральные логарифмы и натуральные (а не двоичные) единицы информации, так как это упрощает проводимые ниже аналитические выкладки.

Декодирующая система для блочкового кода длины n есть метод сопоставления единственного входного сообщения (целого числа от 1 до M) с каждым возможным выходным словом длины n , т. е. функция, сопоставляющая выходным словам длины n целые числа от 1 до M . *Вероятность ошибки* для кода, когда M входных сообщений выбираются каждое с вероятностью $1/M$, есть вероятность того, что шум и декодирующая система приведут ко входному сообщению, отличному от того, которое в действительности было передано.

Если имеются два канала, то двумя естественными способами из них возможно составить один канал, который назовем соответственно суммой и произведением двух первоначальных каналов. *Сумма* двух каналов есть канал, сформированный из входов каждого из двух данных каналов с присущими им вероятностями перехода к множеству выходных букв. При этом множество выходных букв



Р и с. 1.

состоит из логической суммы двух первоначальных выходных алфавитов. Таким образом, суммарный канал определяется матрицей переходов, которая может быть составлена, если расположить матрицу одного канала ниже и правее матрицы другого канала и заполнить два оставшихся прямоугольника нулями. Если $\|p_i(j)\|$ и $\|p_{i'}(j')\|$ являются индивидуальными матрицами, то сумма имеет следующую матрицу:

$$\begin{bmatrix} p_1(1) & \dots & p_1(r) & 0 & \dots & 0 \\ \vdots & & \vdots & \vdots & & \vdots \\ p_t(1) & \dots & p_t(r) & 0 & \dots & 0 \\ 0 & \dots & 0 & p_{i'}(1) & \dots & p_{i'}(r') \\ \vdots & & \vdots & \vdots & & \vdots \\ 0 & \dots & 0 & p_{i'}(1) & \dots & p_{i'}(r') \end{bmatrix}$$

Произведение двух каналов есть канал, входной алфавит которого состоит из всех упорядоченных пар (i, i') , где i — буква из алфавита первого канала, i' — буква из алфавита второго канала, а выходной алфавит состоит из подобной же совокупности упорядоченных пар букв, взятых из обоих индивидуальных выходных алфавитов. У такого канала вероятности перехода из (i, i') в (j, j') равны $p_i(j) p_{i'}(j')$.

Сумма каналов физически соответствует ситуации, в которой может быть использован каждый из двух каналов (но не оба вместе), и выбор делается заново для каждой передаваемой буквы. Произведение каналов соответствует ситуации, в которой оба канала используются вместе в каждую единицу времени. Интересно отметить, что умножение и сложение каналов ассоциативно и коммутативно и что произведение дистрибутивно относительно суммы. Таким образом, для каналов можно разработать некоторого рода алгебру, в которой возможно, например, написать полином $\sum a_n K^n$, где a_n — неотрицательные целые числа, а K — канал. Однако здесь не будут исследоваться алгебраические свойства таких систем.

Пропускная способность канала при нулевой ошибке

Будем говорить, что в дискретном канале две входные буквы являются *смежными*, если существует некоторая выходная буква, которая может быть результатом передачи любой из них. Так i и j смежные, если существует некоторое t , такое, что $p_i(t)$ и $p_j(t)$ не

равны нулю. На рис. 1 буквы a и c являются смежными, тогда как a и d не являются смежными.

Если все входные буквы являются смежными друг с другом, то любой код, содержащий больше одного слова, имеет вероятность ошибки на приемном конце, отличную от нуля. Действительно, вероятность ошибки при декодировании слов удовлетворяет неравенству

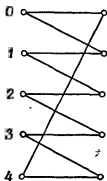
$$P_e \geq \frac{M-1}{M} p_{\min}^n,$$

где $p_{\min}$ — наименьшая (не равная нулю) из вероятностей $p_i(j)$, n — длина кода, а M — число слов в коде. Для того чтобы доказать это, заметим, что любые два слова имеют некоторое общее возможное выходное слово, а именно то слово, которое состоит из последовательности выходных букв, совпадающих при побуквенном сопоставлении двух входных слов. Каждое из двух входных слов имеет вероятность произвести общее выходное слово, по крайней мере равную $p_{\min}^n$. При кодировании каждое из двух входных слов будет появляться в среднем $1/M$ долю раз и будет выдаваться общий выход в среднем $(1/M)p_{\min}^n$ долю раз. Этот выход может быть декодирован лишь единственным образом. Следовательно, по меньшей мере одна из этих ситуаций приводит к ошибке. Эта ошибка, равная $\frac{1}{M} p_{\min}^n$, приписывается данному слову и из оставшихся $M-1$ кодовых слов выбирается другая пара. Ошибка величины $\frac{1}{M} p_{\min}^n$ приписывается подобным же образом одному из этих слов. Эти события оказываются непересекающимися. Продолжая этот процесс далее, получим для общей вероятности ошибки величину, по крайней мере равную $\frac{M-1}{M} p_{\min}^n$.

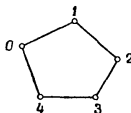
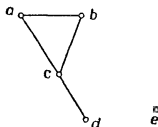
Если условие, что каждая входная буква является смежной со всеми другими буквами, не выполняется, то возможно с положительной скоростью вести передачу при нулевой вероятности ошибки. Наименьшая верхняя граница всех скоростей, которые могут быть достигнуты при нулевой вероятности ошибки, будет называться *пропускной способностью канала при нулевой ошибке* и обозначаться через C_0 . Пусть $M_0(n)$ — наибольшее число слов в коде длины n , никакие два слова которого не являются смежными; тогда C_0 есть верхняя граница чисел $\frac{1}{n} \log M_0(n)$, когда значение n пробегает по всем положительным целым числам.

Можно ожидать, что C_0 будет равна $\log M_0(1)$, т. е., что если выбрать наибольшее возможное множество несмежных букв и составить из них все возможные последовательности длины n , то получится

наилучший код длины n , свободный от ошибок. В общем случае это утверждение не верно, хотя оно и справедливо во многих случаях, особенно при малом числе входных букв. Первое отступление от этого правила происходит в случае пяти входных букв в канале, изображенном на рис. 2. В этом канале возможно выбрать по крайней мере две несмежные буквы, например 0 и 2. Используя составленные из них последовательности 00, 02, 20 и 22, получаем четыре слова в коде длины два. Возможно, однако, сконструировать код длины два из пяти попарно несмежных элементов 00, 12, 24, 31, 43.



Р и с. 2.



Р и с. 3.

Легко проверить, что никакие две из этих комбинаций не являются смежными. Таким образом, C_0 для этого канала, самое меньшее, равна $\frac{1}{2} \log 5$.

В настоящее время не найден метод определения C_0 в случае произвольного дискретного канала. Отыскание этого метода предлагается в качестве интересной нерешенной задачи теории кодирования. Ниже будет получен ряд результатов, которые дают возможность определить C_0 во многих частных случаях. Например, определим C_0 для всех каналов с пятью и меньшим числом входных букв. Единственное исключение составит канал, изображенный на рис. 2 (или каналы, эквивалентные ему по структуре смежности). Также получим некоторые общие неравенства, позволяющие в большинстве случаев весьма точно оценить C_0 .

Можно заметить прежде всего, что величина C_0 зависит только от того, какие входные буквы смежны со всеми остальными входными буквами. Для произвольного канала определим следующим обра-

зом матрицу смежности $\|A_{ij}\|$

$$A_{ij} = \begin{cases} 1, & \text{если входная буква } i \text{ смежна с } j \text{ или если } i=j, \\ 0 & \text{во всех остальных случаях.} \end{cases}$$

Предположим, что два канала имеют (может быть, после перенумерации входных букв одного из них) одну и ту же матрицу смежности. Тогда очевидно, что код, дающий нулевую ошибку в одном из каналов, будет кодом, дающим нулевую ошибку и для другого канала. Поэтому пропускная способность при нулевой ошибке C_0 одного из каналов будет точно так же приложима и к другому каналу.

Структура смежности, отраженная в матрице смежности, может быть также представлена в виде линейного графа. Построим граф с вершинами по числу входных букв и соединим две различные вершины линией или ребром графа, если соответствующие им входные буквы являются смежными. На рис. 3 показаны два примера, соответствующие каналам, изображенным на рис. 1 и 2.

Теорема 1. Пропускная способность дискретного канала без памяти при нулевой ошибке C_0 ограничена неравенствами

$$-\log \min_{P_i} \sum_{i,j} A_{ij} P_i P_j \leq C_0 \leq \min_{p_i(j)} C,$$

$$\sum_i P_i = 1, \quad P_i \geq 0,$$

$$\sum_j p_j(j) = 1, \quad p_i(j) \geq 0.$$

Здесь C — пропускная способность любого канала, имеющего переходные вероятности $p_i(j)$ и матрицу смежности A_{ij} .

Оценка сверху совершенно очевидна. Пропускная способность при нулевой ошибке меньше или равна обычной пропускной способности любого канала с той же самой матрицей смежности, так как первая требует коды с нулевой вероятностью ошибки, тогда как вторая — только коды с вероятностью ошибки, стремящейся к нулю. Минимизируя пропускную способность варьированием $p_i(j)$, найдем наименьшую верхнюю границу, которую можно достичь на основании этих соображений. Так как пропускная способность является непрерывной функцией от $p_i(j)$ в замкнутой области, определяемой $p_i(j) \geq 0$, $\sum_j p_i(j) = 1$, то можно писать $\min$ вместо наибольшей нижней грани.

Стоит заметить, что хотя существует бесконечное число каналов с одной и той же матрицей смежности, при выполнении минимизации необходимо рассматривать только некоторый особый канал. Этот единственный особый канал получается следующим образом из матрицы смежности. Определим выходную букву j с отличными от

нуля $p_i(j)$ и $p_k(j)$, в случае, когда для пары i, k величина $A_{ik} = 1$. Далее, если существуют какие-либо входные буквы, скажем i, k, l , не смежные друг с другом, определим выходную букву, например, m с отличными от нуля $p_i(m)$, $p_k(m)$ и $p_l(m)$. На графе смежности это соответствует полному подграфу с тремя вершинами. После этого подмножествам четырех букв, или полному подграфу с четырьмя вершинами, скажем i, k, l, m , сопоставим выходную букву, которая будет связана с каждой из них, и т. д. Очевидно, что любой канал с той же самой матрицей смежности отличается от только что описанного числом выходных символов для некоторой пары, тройки и т. д. смежных входных букв. В случае если в канале имеется более чем один выходной символ для подмножества смежных входных букв, его пропускная способность лишь уменьшается при идентификации этих символов. Если некоторый канал не содержит символа, скажем для тройки смежных входных букв i, k, l , то это будет специальный случай нашего канонического канала, в котором имеется выходная буква m для этой тройки, в то время как $p_i(m)$, $p_k(m)$ и $p_l(m)$ равны нулю.

Докажем теперь справедливость оценки снизу. Воспользуемся методом построения случайных кодов, основанных на вероятностях P_i для букв. Последние будут выбраны так, чтобы минимизировать квадратичную форму $\sum_{i,j} A_{ij} P_i P_j$. Построим ансамбль кодов, каждый из которых состоит из M слов, содержащих по n букв. Слова в коде выбираются следующим стохастическим методом. Любая буква i в любом слове выбирается независимо от всех остальных букв и с вероятностью P_i . Вычислим теперь в этом ансамбле вероятность того, что какое-либо выбранное слово не является смежным ни с каким другим словом кода. Вероятность того, что первая буква первого слова не смежна с первой буквой второго слова, равна $\sum_{i,j} A_{ij} P_i P_j$, так как здесь случаи смежности суммируются с коэффициентом 1, а случаи несмежности с коэффициентом 0. Вероятность того, что два слова смежны по всем буквам и, следовательно, смежны как слова, равна $(\sum_{i,j} A_{ij} P_i P_j)^n$. Это значит, что вероятность несмежности равна $1 - (\sum_{i,j} A_{ij} P_i P_j)^n$. Вероятность того, что все $M-1$ других слов в коде не смежны данному слову в силу их независимого выбора, равна

$$[1 - (\sum_{i,j} A_{ij} P_i P_j)^n]^{M-1}.$$

Согласно хорошо известному неравенству, эта вероятность больше, чем $1 - (M-1) (\sum_{i,j} A_{ij} P_i P_j)^n$, что в свою очередь больше, чем

$1-M \left(\sum_{i,j} A_{ij} P_i P_j \right)^n$. Если положить $M = (1-\varepsilon)^n \left(\sum_{i,j} A_{ij} P_i P_j \right)^{-n}$, то, выбирая ε достаточно малым, получим скорость, сколь угодно близкую к $-\log \sum_{i,j} A_{ij} P_i P_j$. Далее, если ε уже выбрано, можно, выбирая n достаточно большим, сделать $M \left(\sum_{i,j} A_{ij} P_i P_j \right)^n = (1-\varepsilon)^n$ сколь угодно малым, например меньшим δ . В ансамбле кодов вероятность того, что какое-либо частное слово смежно с любым другим словом в собственном коде, теперь меньше δ . Это значит, что в ансамбле существуют коды, для которых отношение числа таких нежелательных слов к их общему числу в коде меньше или равно δ . В противном случае среднее по ансамблю было бы больше. Выберем такой код и вычеркнем из него слова, обладающие этим нежелательным свойством. Это уменьшит нашу скорость в худшем случае лишь на $\log (1-\delta)^{-1}$. Так как ε и δ произвольно малы, получились свободные от ошибок коды для которых скорость сколь угодно близка к $-\log \min_{P_i} \sum_{i,j} A_{ij} P_i P_j$, как и утверждается в теореме.

В связи с оценкой сверху, данной в теореме 1, следующий далее результат полезен для вычисления минимальной C . Он интересен сам по себе и, кроме того, окажется полезным в дальнейшем при исследовании каналов, имеющих линию обратной связи.

Т е о р е м а 2. Для дискретного канала без памяти с переходными вероятностями $p_i(j)$ и вероятностями входных букв P_i следующие три утверждения являются эквивалентными.

1) Скорость передачи

$$R = \sum_{i,j} P_i p_i(j) \log \frac{p_i(j)}{\sum_k P_k p_k(j)}$$

стационарна при вариации по всем не равным нулю P_i , связанным условием $\sum_i P_i = 1$, и при вариации по таким $p_i(j)$, для которых $P_i p_i(j) \geq 0$ и $\sum_j p_i(j) = 1$.

2) Взаимная информация между парой вход — выход $I_{ij} = \log (p_i(j) / \sum_k P_k p_k(j))$ является постоянной $I_{ij} = I$ для всех пар i, j с неравными нулю вероятностями (т. е. для пар, у которых $P_i p_i(j) > 0$).

3) Функция $p_i(j) = r_j$ является функцией j для всех i , для которых $P_i p_i(j) > 0$, а сумма $\sum_{i \in S_j} P_i = h$ есть постоянная, не

зависящая от j , где S_j — множество входных букв, которые могут быть приняты как выходная буква j с ненулевой вероятностью. При этом $I = \log h^{-1}$.

Величины $p_i(j)$ и P_i , соответствующие максимуму и минимуму пропускной способности при вариации $p_i(j)$ (при вариации, однако, предполагается, что все $p_i(j)$, равные нулю, остаются равными нулю), удовлетворяют всем трем утверждениям теоремы.

Сначала покажем, что эквивалентными являются утверждения 1) и 2), а затем покажем эквивалентность утверждений 2) и 3).

В (конечной) области, определяемой соотношениями $\sum_i P_i = 1$, $P_i \geq 0$, $\sum_j p_i(j) = 1$ и $p_i(j) \geq 0$, функция R является ограниченной и непрерывной функцией своих аргументов P_i и $p_i(j)$. Функция R обладает конечными частными производными по любой $p_i(j) > 0$. В самом деле, как легко подсчитать

$$\frac{\partial R}{\partial p_i(j)} = P_i \log \frac{p_i(j)}{\sum_k P_k p_k(j)}.$$

Необходимое и достаточное условие того, что R постоянна при малых изменениях неотрицательной вероятности $p_i(j)$, подчиняющейся приведенным условиям, состоит в том, чтобы

$$\frac{\partial R}{\partial p_i(j)} = \frac{\partial R}{\partial p_i(k)}$$

при всех i, j и k , для которых P_i , $p_i(j)$ и $p_i(k)$ не равны нулю. Это условие требует, чтобы

$$P_i \log p_i(j) / \sum_m P_m p_m(j) = P_i \log p_i(k) / \sum_m P_m p_m(k).$$

Пусть $Q_j = \sum_m P_m p_m(j)$ — вероятность выходной буквы j , тогда последнее равенство эквивалентно равенству

$$\frac{p_i(j)}{Q_j} = \frac{p_i(k)}{Q_k}.$$

Иначе говоря, $p_i(j)/Q_j$ не зависит от j и является функцией i тогда, когда $P_i > 0$ и $p_i(j) > 0$. Эту функцию от i обозначим через α_i . Таким образом,

$$p_i(j) = \alpha_i Q_j,$$

если равенство $P_i p_i(j) = 0$ не имеет места.

Взяв теперь частную производную R по P_i , получим

$$\frac{\partial R}{\partial P_i} = \sum_j p_i(j) \log \frac{p_i(j)}{Q_i} - 1.$$

Для того чтобы функция была стационарной при условии $\sum_i P_i = 1$, требуется, чтобы $\frac{\partial R}{\partial P_i} = \frac{\partial R}{\partial P_h}$. Таким образом,

$$\sum_j p_i(j) \log \frac{p_i(j)}{Q_j} = \sum_j p_h(j) \log \frac{p_h(j)}{Q_j}.$$

Так как для $P_i p_i(j) > 0$, имеем $p_i(j)/Q_j = \alpha_i$, то последнее равенство переходит в

$$\sum_j p_i(j) \log \alpha_i = \sum_j p_h(j) \log \alpha_h, \\ \log \alpha_i = \log \alpha_h.$$

Таким образом, α_i не зависит от i и может быть записано как α . Следовательно,

$$\frac{p_i(j)}{Q_j} = \alpha, \\ \log \frac{p_i(j)}{Q_j} = \log \alpha = I$$

всегда, когда $P_i p_i(j) > 0$.

Обратный результат легко получается с помощью обращения приведенного выше доказательства. Если

$$\log \frac{p_i(j)}{Q_j} = I,$$

то с помощью простой подстановки в формулу для $\partial R / \partial P_i$, получим $\partial R / \partial P_i = I - 1$. Следовательно, R стационарно при вариации P_i , подчиняющихся условию $\sum P_i = 1$. Далее, $\partial R / \partial p_i(j) = P_i I = \partial R / \partial p_i(k)$ и, следовательно, вариация R обращается в нуль при $\sum_j p_i(j) = 1$.

Теперь докажем, что из 2) вытекает утверждение 3). Предположим, что $\log \frac{p_i(j)}{Q_j} = I$ всегда, когда $P_i p_i(j) > 0$. Тогда $p_i(j) = e^I Q_j$ является функцией j при тех же самых условиях. Также, если $q_j(i)$ является условной вероятностью i при заданном j , то

$$\frac{Q_j q_j(i)}{P_i Q_j} = e^I, \\ q_j(i) = e^I P_i, \\ 1 = \sum_{i \in S_j} q_j(i) = e^I \sum_{i \in S_j} P_i.$$

Чтобы доказать, что из утверждения 3) вытекает утверждение 2), предположим, что

$$p_i(j) = r_j,$$

когда $P_i p_i(j) > 0$. Тогда

$$\frac{P_i p_i(j)}{P_i Q_j} = \frac{r_j}{Q_j} = \lambda_j = \frac{Q_j q_j(i)}{P_i Q_j} = \frac{q_j(i)}{P_i}.$$

Суммируя теперь равенства $P_i \lambda_j = q_j(i)$ по $i \in S_j$ и используя предположение из утверждения 3), состоящее в том, что $\sum_{S_j} P_i = h$,

получим

$$h \lambda_j = 1,$$

так что λ_j равно h^{-1} и не зависит от j . Следовательно, $I_{ij} = I = \log h^{-1}$.

Последнее утверждение теоремы, касающееся минимальной и максимальной пропускной способности при вариации по $p_i(j)$, следует из того факта, что R в этих точках должна быть неизменна при вариации по всем ненулевым P_i и $p_i(j)$. Поэтому соответствующие P_i и $p_i(j)$ удовлетворяют утверждению 1) теоремы.

В простых каналах для оценки C_0 обычно более удобно применить какой-нибудь специальный прием вместо данных в теореме 1 неравенств, которые включают в себя процесс максимизации и минимизации. Наиболее примитивная оценка снизу, как упоминалось ранее, получается просто с помощью отыскания логарифма максимального числа несмежных входных букв.

Весьма полезный аппарат для определения C_0 , который применим во многих случаях, может быть описан, если использовать понятие *отображения, понижающего смежность*. Под этим понимается преобразование некоторых букв в другие буквы, $i \rightarrow \alpha(i)$, обладающее тем свойством, что если i и j несмежны в канале (или на графе), то $\alpha(i)$ и $\alpha(j)$ также не будут смежными. Если имеется код, дающий нулевую ошибку, то можно применить такое отображение буква за буквой ко всему этому коду и получить некоторый новый код, который будет также безошибочным, ибо отображение не добавит смежности между буквами.

Теорема 3. *Если все входные буквы i могут быть преобразованы при помощи отображения $i \rightarrow \alpha(i)$, понижающего смежность, в подмножество попарно несмежных букв, то пропускная способность C_0 канала при нулевой ошибке равна логарифму числа букв этого подмножества.*

Это верно, во-первых, потому, что, если составить всевозможные последовательности из этих букв, получится код, дающий нулевую ошибку при этой скорости передачи. И, во-вторых, любой код для канала, дающий нулевую ошибку, может быть отображен в код, использующий только эти буквы и содержащий, следовательно, по меньшей мере e^{nC_0} несмежных слов.

Пропускная способность при нулевой ошибке или, более точно, эквивалентные числа входных букв для всех графов смежности

вплоть до графов, содержащих пять вершин, показаны на рис. 4. Все они легко могут быть найдены методом теоремы 3. Исключение составляет лишь ранее упоминавшийся канал, изображенный на рис. 2, относительно которого известно только то, что его пропускная способность при нулевой ошибке находится в пределах

$$\frac{1}{2} \log 5 \leq C_0 \leq \log \frac{5}{2}.$$

Нами были рассмотрены все графы с шестью вершинами. Их пропускные способности также могут быть найдены с помощью этой теоремы, за исключением четырех случаев. Эти четыре случая можно изложить в терминах пропускной способности, соответствующей рис. 2, так что этот случай является единственной существенно не разрешенной проблемой вплоть до графов с семью вершинами. Графы с семью вершинами полностью исследованы не были, однако в этом случае возникает по меньшей мере одна новая ситуация, являющаяся аналогом изображенной на рис. 2, но с семью входными буквами.

В качестве примера, иллюстрирующего то, как методом отображения, понижающего смежность, были вычислены величины N_0 , несколько графов на рис. 4 снабжены обозначениями, указывающими на соответствующее отображение. Правило состоит в следующем. Все узлы графа, обозначенные через α , отображаются в узел α , а узел α отображается сам в себя. Все узлы, обозначенные через b и β , отображаются в узел β . Все узлы, обозначенные через c и γ , отображаются в узел γ . Легко проверить, что указанные отображения не создают никаких новых смежностей и что точки α , β и γ не являются смежными.

Пропускная способность C_0 для суммы и произведения каналов

Теорема 4. Пусть имеются два канала без памяти с пропускными способностями при нулевой ошибке $C'_0 = \log A$ и $C''_0 = \log B$. Тогда пропускная способность их суммы при нулевой ошибке больше или равна $\log(A + B)$, а их произведение имеет пропускную способность при нулевой ошибке, большую или равную $C'_0 + C''_0$. В случае когда граф для одного из этих двух каналов может быть сведен к несмежным точкам при помощи метода отображений (теорема 3), эти неравенства могут быть заменены на равенства.

Ясно, что в случае произведения пропускная способность при нулевой ошибке по меньшей мере равна $C'_0 + C''_0$, ибо можно составить код произведения из двух кодов, имеющих скорости, близкие к C'_0 и C''_0 . Если эти коды имеют неравную длину, то для нового кода используем наименьшее общее кратное индивидуальных длин и составим всевозможные последовательности кодовых слов

каждого из кодов вплоть до этой длины. Для того чтобы доказать равенство в случае, когда один из графов (например, для первого канала) может быть отображен в A несмежных точек, предположим, что имеется код для произведения каналов. Буквы для кода произведения являются, конечно, упорядоченными парами букв, соответствующих первоначальным каналам. Заменим первую букву в каждой паре во всех кодовых словах на букву, соответствующую ее преобразованию с помощью метода отображений. Это уменьшит или сохранит смежность между словами в коде. Разделим теперь кодовые слова на A^u подмножеств, согласно последовательностям первых букв в упорядоченных парах. Каждое из этих подмножеств может содержать по крайней мере B^u членов, так как это является наибольшим возможным числом кодов этой же длины для второго канала. Таким образом, в общем, код, дающий желаемый результат, состоит по крайней мере из $A^u B^u$ слов.

Теперь покажем, как в случае суммы двух каналов из заданных двух кодов для двух каналов сконструировать один код для суммарного канала с эквивалентным числом букв, равным $A^{1-\delta} + B^{1-\delta}$, где δ — произвольно малая величина, а A и B — эквивалентные числа букв для того и другого кодов. Пусть коды имеют длины n_1 и n_2 . Новый код будет иметь длину n , где n — наименьшее целое число, большее n_1/δ и n_2/δ . Следующим образом составим теперь коды для первого и для второго каналов для всех длин k от нуля до n . Пусть k равно $an_1 + b$, где a и b — целые числа и $b < n_1$. Составим всевозможные последовательности слов из данного кода для первого канала и заполним все оставшиеся места для b букв произвольно, скажем, первой буквой кодового алфавита. Получим по крайней мере $A^{k-\delta n}$ различных слов длины k , из которых ни одно слово не является смежным ни с каким другим. Тем же самым способом составим коды для второго канала и получим $B^{k-\delta n}$ слов длины k в этом коде. Объединим теперь всеми C_n^k возможными способами код длины k для первого канала с кодом длины $n-k$ для второго канала и произведем это для каждого значения k . При этом возникнет некоторый код длиной в n букв, содержащий по крайней

мере $\sum_{k=0}^n C_n^k A^{k-\delta n} B^{n-k-\delta n} = (AB)^{-\delta n} (A+B)^n$ различных слов.

Легко увидеть, что слова являются попарно несмежными. Соответствующая коду скорость по крайней мере равна $\log(A+B) - \delta \log AB$. В силу того что δ — произвольно малая величина, можно достигнуть скорости, сколь угодно близкой к $\log(A+B)$.

Чтобы показать, что при сведении с помощью отображения одного из графов к несмежным точкам невозможно превысить скорость, соответствующую числу букв $A+B$, рассмотрим какой-либо заданный код длины n для суммарного канала. Слова в нем состоят

из последовательностей букв, каждая из букв в которых соответствует тому или другому каналу. Слова могут быть подразделены на классы в соответствии со способом расстановки букв, принадлежащих тому или другому каналу. Существует всего 2^n таких классов. Среди них имеются C_n^k классов, в которых в точности k букв относятся к первому каналу и $n - k$ ко второму каналу. Рассмотрим теперь некоторый частный класс слов такого типа. Заменяем буквы из алфавита первого канала соответствующими несмежными буквами. Это не повлияет на соотношения смежности между словами в коде. Теперь, как и в случае произведения, расклассифицируем кодовые слова в соответствии с последовательностью букв, включенных в слова для первого канала. Это создаст по крайней мере A^k подмножеств. Каждое из этих подмножеств содержит самое большее B^{n-k} членов, так как это есть наибольшее возможное число несмежных слов длины $n - k$ для второго канала. В общем, если просуммировать по всем величинам k и принять во внимание C_n^k классов для каждого k , то получится, что существуют самое большее $\sum_k C_n^k A^k B^{n-k} = (A + B)^n$ слов в коде для суммарного канала. Это показывает справедливость ожидаемого результата.

Теорема 4 аналогична, конечно, известному положению, относящемуся к обычной пропускной способности C : пропускная способность произведения каналов равна сумме обычных пропускных способностей; эквивалентное число букв суммарного канала равно сумме эквивалентных чисел букв складываемых каналов. Не будучи в состоянии доказать это, мы предполагаем, однако, что равенства в теореме 4 имеют место всегда, а не только при указанных условиях. Теперь получим оценку снизу для вероятности ошибки, в случае когда скорость передачи превышает C_0 .

Теорема 5. При любом коде длины n и скорости $R > C_0$, $C_0 > 0$, вероятность ошибки P_e удовлетворяет неравенству $P_e \geq (1 - e^{-n(C_0 - R)}) p_{\min}^n$, где $p_{\min}$ — минимальная не равная нулю вероятность $p_i^{(j)}$.

Из определения C_0 следует, что существует не более чем e^{nC_0} несмежных слов длины n . Поэтому при $R > C_0$ среди e^{nR} слов должны существовать смежные пары. Смежная пара имеет общее выходное слово, которое может появиться по меньшей мере с вероятностью $p_{\min}^n$. Это выходное слово не может быть декодировано сразу в оба входных слова. Поэтому при декодировании, по крайней мере в одном случае, должна происходить ошибка, когда на выходе появляется общее слово. Это дает вклад в вероятность ошибки P_e , по крайней мере равный $e^{-nR} p_{\min}^n$. Исключим теперь это слово из рассмотрения и применим то же самое рассуждение к оставшимся

$e^{nR} - 1$ словам кода. Это даст другую смежную пару и следующий вклад в ошибку, равный по крайней мере $e^{-nR} p_{\min}^n$. Этот процесс может быть продолжен до тех пор, пока число оставшихся кодовых точек не будет равно в точности e^{nC_0} . К этому моменту вычисленная вероятность ошибки должна быть по крайней мере равна $(e^{nR} - e^{nC_0}) e^{-nR} p_{\min}^n = (1 - e^{n(C_0 - R)}) p_{\min}^n$.

Каналы с обратной связью

Рассмотрим теперь соответствующую задачу для каналов с полной обратной связью. Под этим будем понимать то, что существует «обратный» канал, безошибочно посылающий принятые на самом деле буквы из точки приема обратно на передающий конец. Предполагается, что эта информация принимается на передающем конце до того момента, когда должна быть передана следующая буква. Поэтому при желании она может быть использована при выборе следующей подлежащей передаче буквы.

Интересно, что в канале без памяти обычная пропускная способность в прямом направлении равняется одной и той же величине в присутствии и в отсутствии обратной связи. Это будет показано в теореме 6. Иначе обстоит дело с пропускной способностью при нулевой ошибке, которая в некоторых случаях может быть больше в присутствии обратной связи, чем при ее отсутствии. Например, в канале, изображенном на рис. 5, $C_0 = \log 2$. Однако, как будет показано на основании теоремы 7, пропускная способность при нулевой ошибке и обратной связи $C_{0F} = \log 2,5$.

Сначала определим код из блоков длины n для системы с обратной связью. Это означает, что на передающем конце имеется устройство с двумя входами или математически — функция двух аргументов. Один аргумент — это сообщение, которое должно быть послано, другой — принятые ранее слова, поступившие по каналу обратной связи. Значение функции представляет собой букву, которую следует передать. Таким образом, функция может быть представлена как $x_{j+1} = f(k, v_j)$, где x_{j+1} есть $(j+1)$ -ая буква, переданная в блоке, k — индекс, принимающий значения от 1 до M и выделяющий специфику переданного сообщения, а v_j — принятое слово длины j . Так j изменяется в пределах от 0 до $n-1$; это значит, что v_j принимает значения всех принятых слов этой длины.

Передача осуществляется следующим образом. Если подлежит передаче сообщение m_k , то f вычисляется как $f(k, -)$, где знак минус

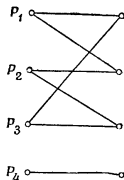


Рис. 5.

означает «слово отсутствует», и посылается первая буква. Если канал обратной связи в качестве первой принятой буквы приносит, например, букву α , то следующей передаваемой буквой будет $f(k, \alpha)$. Если далее принимается β , то следующей переданной буквой будет $f(k, \alpha\beta)$ и т. д.

Теорема 6. *Пропускная способность в прямом направлении дискретного канала без памяти при наличии обратной связи равна обычной пропускной способности C (в отсутствие обратной связи). Среднее приращение взаимной информации I_{vm} принятой последовательности v относительно сообщения m на одну букву текста не превосходит C .*

Пусть v — последовательность, принятая к моменту окончания блока, m — сообщение, x — следующая переданная буква и y — следующая принятая буква. Все они являются случайными величинами, в том числе и x , которая является функцией от m и v . Эта функция фактически является функцией, которая определяется методом кодирования при наличии обратной связи, на основании которого следующая передаваемая буква x определяется сообщением m и информацией v относительно ранее принятых сигналов, поступающей по каналу обратной связи. То, что канал является каналом без памяти, означает, что последующая операция не зависит от предыдущей и, в частности, что $\bar{P}R|y/x| = \bar{P}R|y/x, v|$.

Для среднего приращения взаимной информации, вызываемого парой x, y , при условии, что было принято некоторое заданное v , имеем (усредняя по сообщениям m и последующим принятым буквам y при некотором заданном v)

$$\begin{aligned} \bar{\Delta I} = \overline{I_{m,vy}} - \overline{I_{m,v}} &= \sum_{y,m} Pr|y, m/v| \log \frac{Pr|v, y, m|}{Pr|v, y| Pr|m|} - \\ &- \sum_m Pr|m/v| \log \frac{Pr|v, m|}{Pr|v| Pr|m|}. \end{aligned}$$

Так как $Pr|m/v| = \sum_y Pr|y, m/v|$, вторая сумма может быть записана в виде

$$\sum_{y,m} Pr|y, m/v| \log \frac{Pr|v, m|}{Pr|v| Pr|m|}.$$

Объединяя обе суммы, получим теперь, что

$$\begin{aligned} \bar{\Delta I} &= \sum_{y,m} Pr|y, m/v| \log \frac{Pr|v, y, m| Pr|v|}{Pr|v, m| Pr|v, y|} = \\ &= \sum_{y,m} Pr|y, m/v| \log \frac{Pr|y/v, m| Pr|v|}{Pr|v, y|}. \end{aligned}$$

Помня, что x является функцией m и v , представим суммирование по m как суммирование сначала по тем m , которые дают в результате одно и то же x (при заданном v), а затем по различным x . При первом суммировании величина $Pr|y/v, m|$ является постоянной, равной $Pr|y/x|$, а коэффициенты при логарифме дают в сумме $Pr|x, y/v|$. В результате можно написать

$$\Delta I = \sum_{x, v} Pr|x, y/v| \log \frac{Pr|y/x|}{Pr|y/v|}.$$

Рассмотрим теперь скорость передачи в канале (в обычном смысле при отсутствии обратной связи) так, как будто бы значениям x приписаны вероятности $q(x) = Pr|x/v|$. При этом вероятности $r(x, y)$ для пар x, y и вероятности $w(y)$ просто для y были бы равны

$$r(x, y) = q(x) Pr|y/x| = Pr|x/v| Pr|y/x| = Pr|x, y/v|,$$

$$w(y) = \sum_x r(x, y) = \sum_x Pr|x, y/v| = Pr|y/v|.$$

Скорость, следовательно, равнялась бы

$$R = \sum_{x, v} r(x, y) \log \frac{Pr|y/x|}{w(y)} = \sum_{x, v} Pr|x, y/v| \log \frac{Pr|y, x|}{Pr|y/v|} = \Delta I.$$

Так как $R \leq C$, где C — пропускная способность канала (C является максимально возможным значением R по всем $q(x)$), заключаем, что

$$\Delta I \leq C.$$

Так как среднее изменение I на одну букву не превосходит C , среднее изменение на n букв не превосходит nC . Следовательно, в блоковом коде длины n при скорости на входе, равной R , $R > C$, ненадежность при окончании блока будет по крайней мере $R - C$, т. е. точно такая же, как в случае канала без обратной связи. Иначе говоря, невозможно достигнуть нулевой ненадежности (или, что легко следует отсюда, нулевой вероятности ошибки) в случае, когда скорость превышает пропускную способность канала. Это, конечно, возможно сделать при скоростях, меньших C , ибо, наверняка, все, что может быть сделано без обратной связи, может быть сделано и с обратной связью.

Интересно, что первое утверждение теоремы 6 может быть легко обобщено на каналы с памятью при условии, что они имеют такую структуру, что внутреннее состояние канала может быть вычислено на передающем конце по начальному состоянию и последовательности посланных букв. Если этот случай не имеет места, заключение

теоремы не будет всегда справедливым, т. е. существуют каналы более сложной структуры, для которых при наличии обратной связи пропускная способность в прямом направлении превосходит ту, которая имеет место при отсутствии обратной связи. Здесь, однако, не будут приводиться подробности этих обобщений.

Обращаясь теперь вновь к задачам передачи с нулевой ошибкой, дадим определение пропускной способности при нулевой ошибке C_{0F} для канала с обратной связью следующим очевидным образом: пропускная способность есть верхняя граница скоростей для блоковых кодов, которые не приводят к ошибкам. Следующая теорема решает задачу вычисления C_{0F} для канала без памяти при наличии обратной связи и указывает, как быстро при увеличении длины блока n можно достигнуть скорости C_{0F} .

Теорема 7. *Пропускная способность дискретного канала без памяти при нулевой ошибке и при наличии полной обратной связи, воспроизводящей принятые буквы на передающем конце, равна нулю, если все пары входных букв являются смежными. В противном случае $C_{0F} = \log P_0^{-1}$, где*

$$P_0 = \min_{P_i} \max_i \sum_{i \in S_j} P_i,$$

P_i — заданная вероятность входной буквы i ($\sum_i P_i = 1$), а S_j — множество входных букв, переходящих в выходную букву j с положительной вероятностью. Для такого канала с обратной связью при передаче со скоростью $R \geq C_{0F}(1 - \frac{2}{n} \log_2 2t)$, где t — число входных букв, можно найти не дающий ошибок код из блоков длины n .

Появившаяся в этой теореме величина P_0 имеет следующий смысл. При любых заранее заданных вероятностях входных букв P_i для каждой выходной буквы j можно вычислить совокупную вероятность всех входных букв, переходящих (с положительной вероятностью) в букву j . Это будет $\sum_{i \in S_j} P_i$. Выходные буквы, для которых сумма будет очень большой, являются «плохими» с той точки зрения, что, когда они приняты, существует большая неопределенность, относящаяся к переходам. Чтобы получить P_0 , следует выбрать P_i так, чтобы «худшее» в указанном смысле выходное слово было бы возможно «лучшим».

Покажем сначала, что $C_{0F} = 0$, если все буквы являются взаимно смежными. В самом деле, какова бы ни была система кодирования, любые два сообщения, например m_1 и m_2 , могут с положительной вероятностью дать при декодировании одну и ту же выходную последовательность. А именно первые переданные буквы, соответ-

ствующие m_1 и m_2 , имеют общую возможную выходную букву. Предполагая, что такая буква появилась на выходе, обратимся к следующим передаваемым для m_1 и m_2 буквам в той же системе кодирования. Они также имеют общую возможную выходную букву. Продолжая этот процесс, установим некоторое выходное слово, которое может быть результатом приема любой из m_1 и m_2 , и поэтому сообщения не могут быть однозначно распознаны.

Рассмотрим теперь такой случай, при котором не все пары являются смежными. Сначала, введя в рассмотрение блочный код длины n , докажем, что при кодировании, не дающем ошибок, нельзя превзойти скорости $\log P_0^{-1}$. Для $n = 0$ это утверждение, очевидно, справедливо. Пусть доказываемая по индукции гипотеза будет состоять в том, что не существует блочного кода длины $n-1$, передающего со скоростью, превышающей $\log P_0^{-1}$, или, иначе говоря, нельзя однозначно различить более чем $e^{(n-1) \log P_0^{-1}} = P_0^{-(n-1)}$ различных сообщений. Предположим теперь (в противоположность ожидаемому результату), что имеется блочный код длины n , различающий M сообщений при $M > P_0^{-n}$. Первая передаваемая в коде буква делит эти M сообщений между входными буквами канала. Пусть F_i — доля сообщений, приписанных к букве i (т. е. таких сообщений, для которых i является первой передаваемой буквой). Эти F_i теперь подобны вероятностям, заданным для различных букв, и поэтому, согласно определению P_0 , существует некоторая выходная буква, скажем буква k , такая, что $\sum_{i \in S_j} F_i \geq P_0$. Рассмотрим множество

сообщений, для которых первая передаваемая буква принадлежит S_k . Число сообщений в этом множестве равно по меньшей мере $P_0 M$. Каждое из них может дать выходную букву k в качестве первой принимаемой буквы. Тогда, когда это случится, существуют еще $n-1$ букв, подлежащих передаче, и так как $M > P_0^{-n}$, то имеем $P_0 M > P_0^{-(n-1)}$. Таким образом, существует не дающий ошибок блочный код длины $n-1$, дающий скорость передачи, большую чем $\log P_0^{-1}$, что противоречит предположению индукции. Отметим, что кодирующая функция для кода длины $n-1$ определяется формально из первоначальной кодирующей функции с помощью фиксирования k в качестве применяемой буквы.

Теперь требуется показать, что при кодировании, не дающем ошибок, можно на самом деле достигнуть скорости, сколь угодно близкой к $\log P_0^{-1}$. Пусть P_i — совокупность вероятностей, которые, будучи приписанными входным буквам, дадут P_0 для $\min_i \max_j \sum_{i \in S_j} P_i$. Тогда общая схема кода будет заключаться в разделении M первоначальных сообщений на t различных групп в соответствии с первой передаваемой буквой. Число сообщений в этих группах будет приблизительно пропорционально $P_1, P_2, \dots, P_t$.

Первая передаваемая буква будет соответствовать группе, содержащей сообщение, которое должно быть передано. Какая бы буква ни была принята, число сообщений, совместимых с этой принятой буквой, будет приблизительно $P_0 M$. Это подмножество возможных сообщений известно как на приемном, так и на передающем конце (после того, как принятая буква посылается обратно к передатчику).

Кодирующая система далее разделяет это подмножество сообщений на t групп, снова приблизительно пропорционально вероятностям P_i . Вторая передаваемая буква соответствует группе, содержащей существующее сообщение. Какая бы буква ни была принята, число сообщений, совместимых с обеими принятыми буквами, равняется теперь примерно $P_0^2 M$.

Этот процесс продолжается до тех пор, пока лишь незначительное число сообщений (меньшее t^2) не станет совместимым со всеми принятыми буквами. Неопределенность между ними может быть снята, если использовать пары несмежных букв в простом двоичном коде. Код, построенный таким образом, будет кодом с нулевой ошибкой для рассматриваемого канала.

Теперь более тщательно оценим аппроксимацию, которая имеет место при подразделении сообщений на t групп. Покажем, что для любого M и любой совокупности P_i ($\sum_i P_i = 1$) можно разделить M сообщений на группы по $m_1, m_2, \dots, m_t$ сообщений, такие, что $m_i = 0$ всегда, когда $P_i = 0$ и

$$\left| \frac{m_i}{M} - P_i \right| \leq \frac{1}{M}, \quad i = 1, \dots, t.$$

Без ограничения общности предположим, что $P_1, P_2, \dots, P_s$ не равны нулю. Выбираем m_1 так, чтобы это было наибольшее целое число, такое, что $m_1/M \leq P_1$. Пусть $P_1 - m_1/M = \delta_1$. Ясно, что $|\delta_1| \leq 1/M$. Выберем далее m_2 так, чтобы это было наименьшее целое число, такое, что $m_2/M \geq P_2$. Положим $P_2 - m_2/M = \delta_2$. Имеем $|\delta_2| \leq 1/M$. Точно так же $|\delta_1 + \delta_2| \leq 1/M$, так как δ_1 и δ_2 противоположны по знаку и каждая из них по абсолютной величине меньше $1/M$. Далее выберем m_3 так, чтобы m_3/M было приблизительно, с точностью $1/M$, равно P_3 . Если $\delta_1 + \delta_2 \geq 0$, то m_3/M выбирается меньшим или равным P_3 . Если $\delta_1 + \delta_2 < 0$, то m_3/M выбирается большим или равным P_3 . Таким образом, снова $P_3 - m_3/M = \delta_3 \leq 1/M$ и $|\delta_1 + \delta_2 + \delta_3| \leq 1/M$. Продолжая этот процесс до P_{s-1} включительно, получим приближения для $P_1, P_2, \dots, P_{s-1}$, обладающие тем свойством, что $|\delta_1 + \delta_2 + \dots + \delta_{s-1}| \leq 1/M$, или $|M(P_1 + P_2 + \dots + P_{s-1}) - (m_1 + m_2 + \dots + m_{s-1})| \leq 1$. Если теперь определить m_s как $M - \sum_{i=1}^{s-1} m_i$, это неравенство может

быть записано в виде $|M(1 - P_s) - (M - m_s)| \leq 1$. Отсюда $|m_s/M - P_s| \leq 1/M$. Таким образом, так как все приближения m_i/M лежали в пределах $1/M$ от P_i , а $\sum m_i = M$, цель достигнута.

Возвращаясь теперь к нашей основной задаче, отметим сначала, что если $P_0 = 1$, то $C_{0F} = 0$, и справедливость теоремы тривиальна. Предположим затем, что $P_0 < 1$. Покажем, что $P_0 \leq 1 - 1/t$. Рассмотрим множество входных букв, которые имеют максимальное значение P_i . Этот максимум, конечно, больше или равен среднему значению, равному $1/t$. Можно далее сделать так, чтобы по крайней мере одна из этих входных букв не была связанной с некоторой выходной буквой. Предположим, что этого сделать нельзя. Тогда либо не существует других входных букв, кроме этого множества, что противоречит предположению, что $P_0 \leq (1 - 1/t)$, либо существуют другие входные буквы с меньшими значениями P_i . В этом случае уменьшение P_i для одной из входных букв в максимальном множестве и увеличение соответственно вероятности для какой-либо входной буквы, не связанной со всеми выходными буквами, не увеличит значения P_0 (при любом S_j) и даст входную букву желаемого типа. Рассматривая выходные буквы, с которыми эта входная буква не связана, видим, что $P_0 \leq (1 - 1/t)$.

Предположим теперь, что передача начинается с M сообщений, подразделенных на группы, вероятности которых, согласно описанному ранее, приблизительно пропорциональны P_i . Тогда, если некоторая буква является уже принятой, множество возможных сообщений (совместимых с этой принятой буквой) сводится к тем, которые объединены в группы, соответствующие буквам, соединенным с уже принятой буквой. Каждая выходная буква соединена не более чем с $t - 1$ входными буквами (в противном случае имелось бы равенство $P_0 = 1$). Для любой из образованных групп ошибка в приближении P_i меньше или равна $1/M$. Следовательно, общее относительное число сообщений во всех соединенных группах меньше или равно $P_0 + (t - 1)M$ для любой выходной буквы. Общее число возможных сообщений после приема первой буквы снижается поэтому от M до числа, меньшего или равного $P_0 M + t - 1$.

В системе кодирования, которая была использована, оставшееся возможное подмножество сообщений снова делится между входными буквами так, чтобы аппроксимировать точно таким же образом вероятности P_i . Это подразделение может быть выполнено с использованием того же стандартного приема (в точности такого же, как описанный выше), как в точке приема, так и на передающем конце, в силу того, что обратная связь для обоих из них делает доступным требуемые данные, а именно первую принятую букву.

Вторая передаваемая буква, полученная таким образом, в свою очередь сведет на приемном конце число возможных сообщений

к величине, не превышающей $P_0(P_0M + t - 1) + t - 1$. Тот же самый процесс продолжается для каждой передаваемой буквы. Если оценка сверху числа возможных остающихся сообщений после k букв есть M_k , то $M_{k+1} = P_0M_k + t - 1$. Решением этого разностного уравнения является

$$M_k = AP_0^k + \frac{t-1}{1-P_0}.$$

Это легко может быть проверено подстановкой в разностное уравнение. Для того чтобы удовлетворить начальным условиям $M_0 = M$, требуется, чтобы $A = M - \frac{t-1}{1-P_0}$. Таким образом, решением оказывается выражение:

$$M_k = \left(M - \frac{t-1}{1-P_0}\right)P_0^k + \frac{t-1}{1-P_0} = MP_0^k + \frac{t-1}{1-P_0}(1-P_0^k) \leq MP_0^k + t(t-1),$$

так как выше было показано, что $1 - P_0 \geq 1/t$.

Если описанный процесс выполнить для n_1 шагов, где n_1 — наименьшее целое число не менее d , а d — решение уравнения $MP_0^d = 1$, то число оставшихся сообщений, совместимых с принятой последовательностью, не будет превышать $1 + t(t-1) \leq t^2$ (так как $t > 1$, в противном случае $C_{0F} = 0$). Теперь пары несмежных букв, существование которых предполагается в теореме, могут быть использованы для разрешения неопределенности по этим t^2 (или меньшего числа) сообщений. Это потребует не более чем $1 + \log_2 t^2 = \log_2 2t^2$ дополнительных букв. Таким образом, в общем можно использовать не большую чем $d + 1 + \log_2 2t^2 = d + 1 + \log_2 4t^2 = n$ длину блока. При такой длине блока передается одно сообщение, выбранное из $M = P_0^{-d}$ возможных. Следовательно, скорость передачи при нулевой ошибке, которая будет достигнута, равна

$$R = \frac{1}{n} \log M \geq \frac{d \log P_0^{-1}}{d + \log_2 4t^2} = \left(1 - \frac{1}{n} \log 4t^2\right) \log P_0^{-1} = \left(1 - \frac{1}{n} \log 4t^2\right) C_{0F}.$$

Итак получено сколь угодно близкое приближение к C_{0F} при помощи кодов, не дающих ошибок.

В качестве примера к теореме 7 рассмотрим канал, изображенный на рис. 5. Требуется вычислить P_0 . Легко увидеть, что, составляя минимакс, входящий в формулировку теоремы 7, можно положить $P_1 = P_2 = P_3$, так как, если бы они были неравными, максимум $\sum_{i \in S_j} P_i$ для соответствующих трех выходных букв можно было бы снизить уравнением. Очевидно также, что тогда $P_4 = P_1 + P_2$.

ибо в противном случае сдвиг вероятностей в ту или другую сторону привел бы к понижению максимума. Отсюда можно заключить, что $P_1 = P_2 = P_3 = 1/5$, а $P_4 = 2/5$. Окончательно пропускная способность канала с обратной связью при нулевой ошибке равна $\log P_0^{-1} = \log 5/2$.

Существует тесная связь между минимаксным процессом теоремы 7 и процессом отыскания минимальной пропускной способности канала при вариациях переходных вероятностей $p_i(j)$, подобным тем, которые имели место в теореме 2. Там было отмечено, что при минимуме пропускной способности каждая выходная буква может получиться из входных букв, имеющих одну и ту же общую вероятность. В самом деле, кажется весьма правдоподобным то, что вероятности входных букв, на которых достигается минимум пропускной способности, в точности совпадают с вероятностями, которые решают проблему минимакса в теореме 7, и если это так, то $C_{\min} = \log P_0^{-1}$.

Автор приносит благодарность Питеру Элайесу за первое указание на то, что линия обратной связи может увеличить пропускную способность при нулевой ошибке, а также за ряд советов, которые были полезными при доказательстве теоремы 7.

ГЕОМЕТРИЧЕСКИЙ ПОДХОД К ТЕОРИИ ПРОПУСКНОЙ СПОСОБНОСТИ КАНАЛОВ СВЯЗИ¹⁾

Методам вычисления скорости передачи информации R и пропускной способности C дискретного канала без памяти²⁾, может быть дано геометрическое толкование, которое приводит к новым результатам и новому пониманию свойств этих величин. Наши результаты обобщают интересную статью Мурога³⁾ и в некоторой степени перекрываются с ней, хотя мы исходили из различных соображений. Метод нашего исследования совершенно иной, поскольку нами используется геометрический подход, основанный на результатах теории выпуклых тел⁴⁾ в противоположность алгебраическому подходу, который использовал Мурога.

Пусть канал определен матрицей $\| p_i(j) \|$ вероятностей перехода от буквы i на входе к букве j на выходе ($i = 1, 2, \dots, a$; $j = 1, 2, \dots, b$). Можно рассматривать каждую строку этой матрицы как вектор или точку в $(b - 1)$ -мерном равностороннем симплексе $[(b - 1)$ -мерный аналог отрезка прямой, равностороннего треугольника, тетраэдра и т. д.]. Координатами точки являются ее расстояния от граней, в сумме равные единице $\sum_j p_i(j) = 1$. Они известны под названием барицентрических координат и соответствуют, например, координатам, часто используемым химиками при описаниях сплавов в терминах долей различных компонент.

Таким образом, с входом i связывается точка или вектор A_i . Его компоненты равны вероятностям различных букв на выходе, если используются все входы. Если использованы все входы (с вероятностью P_i для входа i), то вероятности букв на выходе

¹⁾ Shannon C., Geometrische Deutung einiger Ergebnisse bei der Berechnung der Kanalkapazität, *Nachrichtentechnische Zeitschrift*, 10, 1957, 1.

²⁾ Shannon C., A mathematical theory of communication. (Русский перевод см. стр. 243 данного сборника.—Прим. ред.)

³⁾ Muroga S., On the capacity of a discrete channel, *J. Phys. Soc. Japan*, 8, 4 (1953).

⁴⁾ Alexandroff P., Hopf H., *Topologie I*, Berlin, 1935. (Из доступной литературы см. Александров А. Д., Выпуклые многогранники, Гостехиздат, 1950.—Прим. ред.)

даются компонентами векторной суммы

$$\underline{Q} = \sum_i P_i \underline{A}_i.$$

Вектору или точке $\underline{Q}$ симплекса соответствуют также вероятности букв на выходе. Тогда j -я компонента этого вектора равна $\sum_i P_i p_i(j)$. Поскольку P_i неотрицательны и в сумме дают единицу, то точки $\underline{Q}$ лежат в выпуклой оболочке (или барицентрической оболочке) точек $\underline{A}_j$. Более того, любая точка в этой выпуклой оболочке может быть получена при подходящем выборе P_j .

Теперь для удобства обозначений определим *энтропию* точки или вектора из симплекса как энтропию барицентрических координат точки, интерпретированных как вероятности. Таким образом, имеем

$$H(\underline{A}_i) = - \sum_j p_i(j) \log p_i(j), \quad i = 1, 2, \dots, a,$$

и

$$H(\underline{Q}) = - \sum_j \left[\sum_i P_i p_i(j) \log \right] \sum_i P_i p_i(j),$$

где $H(\underline{Q})$ — энтропия распределения полученных символов.

В этих обозначениях скорость передачи R для данной системы вероятностей на входе P_j задается формулой

$$R = H\left(\sum_i P_i \underline{A}_i\right) - \sum_i P_i H(\underline{A}_i) = H(\underline{Q}) - \sum_i P_i H(\underline{A}_i)^1.$$

Функция $H(\underline{Q})$, где $\underline{Q}$ — точка симплекса, является выпуклой кверху функцией. Так, если компоненты $\underline{Q}$ равны x_i , то имеем

$$H = - \sum x_i \log x_i, \quad \frac{\partial H}{\partial x_i} = -(1 + \log x_i),$$

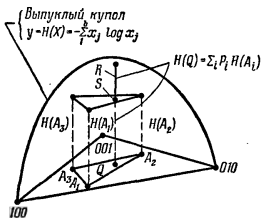
$$H_{ij} = \frac{\partial^2 H}{\partial x_i \partial x_j} = \begin{cases} 0 & , i \neq j, \\ -1/x_i & , i = j. \end{cases}$$

Отсюда $\sum_{i,j} H_{ij} \Delta x_i \Delta x_j = - \sum_i \frac{1}{x_i} (\Delta x_i)^2$ является отрицательно определенной формой. Это справедливо для пространства всех неотрицательных x_i и отсюда, конечно, и для подпространства, в котором $\sum_i x_i = 1$. Следовательно, скорость R , о которой шла речь выше,

¹⁾ Скорость передачи и пропускную способность удобнее определить здесь в битах на букву, а не в битах на символ.

всегда неотрицательна. H строго выпукла (без плоских участков), и R положительна, если только $\underline{A}_i$ не равно $\underline{Q}$ для всех тех i , для которых $P_i > 0$.

Процесс вычисления R может быть легко изображен наглядно в случае двух или трех букв на выходе. Если на выходе имеется три буквы, то представим себе равносторонний треугольник на некоторой основной плоскости. Это будет симплекс, содержащий точки $\underline{A}_j$ и $\underline{Q}$. Сверху этот треугольник покрыт куполообразной поверхностью, как показано на рис. 1. Высота этой поверхности над



Р и с. 1. Пропускная способность канала в случае трех букв на входе и на выходе.

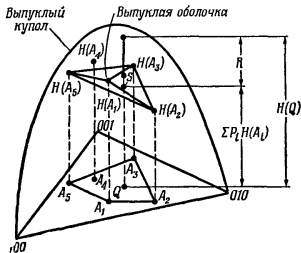
любой точкой $\underline{A}$ равна $H(\underline{A})$. Если на входе имеются три буквы с соответствующими векторами $\underline{A}_1, \underline{A}_2, \underline{A}_3$, то они отвечают трем точкам в треугольнике и трем точкам на куполе, расположенным непосредственно над первыми тремя точками. Каждый выходной вектор $\underline{Q} = \sum_i P_i (\underline{A}_i)$ является точкой треугольника, находящегося на основной плоскости и определенного точками $\underline{A}_1, \underline{A}_2, \underline{A}_3$.

Энтропия $H(\underline{Q})$ равна высоте купола над точкой $\underline{Q}$, а $\sum_i P_i H(\underline{A}_i)$ равна высоте над точкой $\underline{Q}$ плоскости, задаваемой тремя точками купола, расположенными над $\underline{A}_1, \underline{A}_2, \underline{A}_3$. Иными словами, R равно вертикальному отрезку прямой, проходящей через $\underline{Q}$, который отсекается куполом и плоскостью, определяемой этими тремя точками.

Пропускная способность S равна максимуму R . Следовательно, в этом частном случае она равна максимальному расстоянию в вертикальном направлении от купола до внутренней части треугольника, вершины которого расположены на куполе над точками $\underline{A}_1, \underline{A}_2, \underline{A}_3$. Этот максимум, очевидно, достигается в точке касания

плоскости, касательной к куполу и параллельной плоскости того же треугольника при условии, что проекция точки касания на основную плоскость лежит внутри треугольника, заданного точками A_1, A_2, A_3 . В противном случае максимум R достигается на одной из сторон этого треугольника.

Если бы имелось четыре буквы на входе, то они в зависимости от своего расположения определили бы треугольник или четырехугольник на основной плоскости, а точки, расположенные прямо над ними на куполе, определили бы, вообще говоря, тетраэдр.



Р и с. 2. Пропускная способность канала для пяти входных букв и трех выходных букв.

Использование для букв на входе различных распределений вероятности приводит к различным точкам в тетраэдре и к различным значениям $\sum P_i H(A_i)$, вычитаемым при вычислении R . Очевидно, что максимум R был бы достигнут только при таком выборе вероятностей, когда эта вычитаемая часть лежит где-либо на нижней грани тетраэдра. Эти замечания применимы также в случае, когда имеется еще большее количество букв на входе. Если имеется a букв на входе, то они определяют многоугольник с a или меньшим числом сторон на основной плоскости, а точки на куполе, расположенные над вершинами этого многоугольника, образуют симплекс. Любая точка из выпуклой оболочки точек, полученная на куполе, достигается при подходящем выборе P_i ; ей соответствует некоторое значение вычитаемого в формуле для вычисления R . Ясно, что для нахождения максимума R и определения таким образом значения S необходимо рассмотреть только нижнюю часть поверхности выпуклой оболочки (см. рис. 2).

Геометрически очевидно также, что из выпуклости книзу нижней части симплекса и *строгой* выпуклости кверху купола вытекает существование единственной точки, в которой достигается максимум R , и значит, и C . Если бы имелись две такие точки, то значение R в середине отрезка, соединяющего эти две точки, давало бы еще лучшее значение. Это объясняется тем, что вдоль кривой на куполе, соединяющей проекции этих точек, поверхность купола выпукла вверх, а средняя точка нижней соединительной кривой (на нижней грани выпуклой оболочки) не может сдвинуться вверх. Далее, скорость R является строго выпуклой кверху функцией выходного вектора $\underline{Q}$. Оказывается справедливым также тот факт, что скорость R является выпуклой кверху функцией вектора входных вероятностей (этот вектор имеет a барицентрических координат $P_1, P_2, \dots, P_a$ в противоположность b координатам других наших векторов). Это утверждение выполняется потому, что векторы $\underline{Q}'$ и $\underline{Q}''$, отвечающие вероятностям на входе P'_i и P''_i , даются формулами

$$\underline{Q}' = \sum P'_i \underline{A}_i, \quad \underline{Q}'' = \sum P''_i \underline{A}_i.$$

Точкой $\underline{Q}$, отвечающей вектору $\alpha P'_i + \beta P''_i$ (где $\alpha + \beta = 1$ и α и β положительны), является точка $\alpha \underline{Q}' + \beta \underline{Q}''$ и, следовательно, $R \geq \alpha R' + \beta R''$, что и требовалось установить. Равенство может встретиться в случае $\underline{Q}' = \underline{Q}''$, так что в этом случае нельзя говорить о *строго* выпуклой функции.

Из этих последних замечаний вытекает также, что множество S векторов P , для которых скорость R равна пропускной способности C , образует выпуклое множество в своем a -мерном симплексе. Если максимум достигается в двух различных точках, то он также достигается и во всех точках отрезка, соединяющего эти точки. Более того, любой локальный максимум R является также абсолютным максимумом, равным C . Допустим, что это неверно, и соединим отрезком точки, соответствующие локальному и абсолютному максимумам. Значение R должно в силу выпуклости лежать на этой линии или выше ее, но в силу свойств максимума оно должно лежать ниже ее в достаточной близости от локального максимума. Это противоречие доказывает наше утверждение.

Приведенные результаты для геометрической наглядности были описаны для случая трех букв на выходе, но они легко обобщаются на случай n букв на выходе, если использовать хорошо известные результаты теории выпуклых тел.

Еще одно легко выводимое свойство канала состоит в том, что пропускная способность C может быть достигнута при использовании не более чем q букв на входе, где q — ранг матрицы $\|p_i(j)\|$. Это объясняется тем, что $q - 1$ равно размерности множества точек

$\underline{A}_i$. Любая точка на поверхности $(q-1)$ -мерного симплекса лежит на некоторой его грани. Эта грань может быть подразделена на $(q-1)$ -мерные симплексы (если она сама еще не является симплексом). Значит, точка лежит в одном из них. Вершинами симплекса являются q букв на входе, и рассматриваемая точка может быть выражена через них. Теперь легко получается результат Мурога, состоящий в том, что пропускная способность не превышает $\log q$. Действительно, если использованы только q букв, то энтропия входа не может превысить $\log q$, а ненадежность может только уменьшить её значение.

Геометрическая картина дает важную информацию относительно того, какие буквы на входе следует использовать для достижения пропускной способности канала. Если, скажем, вектор $\underline{A}_t$, соответствующий входной букве t , лежит в выпуклой оболочке векторов, соответствующих остальным буквам, то его не нужно использовать. Так, предположим, что $\underline{A}_t = \sum_{i \neq t} a_i \underline{A}_i$, где $\sum_i a_i = 1$, $a_i \geq 0$. Тогда по свойству выпуклости $H(\underline{A}_t) \geq \sum_{i \neq t} a_i H(\underline{A}_i)$. Если при использовании $\underline{A}_i$ с вероятностями P_i получаем скорость

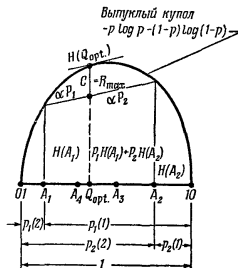
$$R = H\left(\sum_i P_i \underline{A}_i\right) - \sum_i P_i H(\underline{A}_i),$$

то скорость, большая или равная R , может быть получена, если в этой формуле выразить $\underline{A}_t$ через другие $\underline{A}_i$; эта операция не изменит первый член R и уменьшит или не изменит вычитаемую сумму.

В случае наличия только двух букв на выходе ситуация чрезвычайно проста. Каково бы ни было количество букв на входе, для достижения пропускной способности канала нужно использовать только две из них. Эти две буквы должны быть выбраны так, чтобы на них достигались максимум и минимум вероятностей перехода к одной из букв на выходе. Эти буквы, скажем $p_1(2)$ и $p_2(1)$, помещены в одномерном симплексе — отрезке единичной длины и проектируются кверху на кривую H , как показано на рис. 3. Если провести секущую, то пропускная способность будет равна наибольшему расстоянию в вертикальном направлении от секущей до кривой. Вероятности p_1 и p_2 , которые необходимы для достижения этой пропускной способности, пропорциональны расстояниям от этой точки до обоих концов секущей.

В случае трех букв на выходе концы всех векторов, соответствующих буквам на входе, могут быть изображены точками равнобедренного треугольника. Можно рассмотреть многоугольник, ограничивающий эти точки (их выпуклую оболочку вычеркнуть), и точки, внутренние для этого многоугольника (включая и точки на контуре),

требуется рассмотреть нижнюю поверхность симплекса, определенных точками, расположенными на куполе (над оставшимися точками). Эта нижняя поверхность, вообще говоря, будет состоять из треугольников, и наша задача заключается в нахождении вершин, лежащих на этой нижней поверхности. Метод решения этой задачи состоит, например, в рассмотрении отрезка, соединяющего пару вершин, и в последующем выяснении, лежат ли выше или ниже



Р и с. 3. Пропускная способность канала в случае двух букв на выходе.

этого отрезка другие отрезки, проекции которых на основную плоскость пересекают проекцию первого отрезка. Если нет отрезков, лежащих ниже первого, то этот отрезок является ребром нижней поверхности симплекса. Если же некоторый отрезок лежит ниже первого, то он может быть исследован аналогичным способом, и в конце концов будет найдено одно ребро. Это ребро делит проекцию на два меньших многоугольника, и каждый из этих многоугольников должен быть рассмотрен таким же способом. В конце концов начальный многоугольник будет разбит на систему многоугольников, соответствующих поверхности симплекса. Относительно каждого из этих многоугольников нужно затем проверить, лежит ли или нет в этом многоугольнике проекция точки касания плоскости, параллельной ему соответствующей грани, касательной к куполу. Это случится в точности в одном из многоугольников, и тем самым будет определена точка Q , в которой достигается максимум R .

Теперь докажем другое свойство выпуклости для дискретных каналов: докажем, что пропускная способность канала с переход-

ными вероятностями $p_i(j)$ является выпуклой книзу функцией этих вероятностей. Это значит, что пропускная способность C удовлетворяет неравенству, аналогичному неравенству соответствующих переходных вероятностей:

$$r_i(j) = \frac{1}{2} [p_i(j) + q_i(j)],$$

$$C \leq \frac{1}{2} C_1 + \frac{1}{2} C_2,$$

где C_1 является пропускной способностью, отвечающей вероятностям $p_i(j)$, а C_2 — пропускной способностью, отвечающей вероятностям $q_i(j)$.

Чтобы доказать это, допустим, что пропускная способность канала с переходными вероятностями $r_i(j)$ достигается, когда вероятности букв на входе равны P_i . Рассмотрим следующий канал. Он имеет такое же число входных букв, как и данный канал, и удвоенное число выходных букв (которое разобьем на два равных множества $\{j\}$ и $\{j'\}$). Каждой выходной букве соответствуют переходные вероятности $\frac{1}{2} p_i(j)$ и $\frac{1}{2} p_i(j')$. Таким образом, этот канал является каналом, который получился бы при делении пополам всех вероятностей в каналах, отвечающих переходным вероятностям $p_i(j)$ и $p_i(j')$, и при одновременном отождествлении соответствующих входных букв и оставлении выходных букв различными. Заметим, что если соответственные выходные буквы отождествлены, то канал сводится к каналу, отвечающему переходным вероятностям $r_i(j)$. Заметим также, что без такого отождествления этот канал будет работать так же, как канал, который половину времени работает как канал с переходными вероятностями $p_i(j)$, а половину времени как канал с переходными вероятностями $p_i(j')$. Отождествление некоторых выходов всегда уменьшает (или не изменяет) скорость передачи. Пусть этот канал используется с вероятностями P_i для входных символов. Тогда для скоростей передачи может быть написано такое неравенство

$$H(x) - \left[\frac{1}{2} H_{y_1}(x) + \frac{1}{2} H_{y_2}(x) \right] \geq H(x) - H_y(x) = C,$$

где $H_{y_1}(x)$ является условной энтропией x , когда y принадлежит группе $\{j\}$, а $H_{y_2}(x)$ является условной энтропией x , когда y принадлежит группе $\{j'\}$. Разбивая $H(x)$ на две части, чтобы объединить их с $H_{y_1}(x)$ и $H_{y_2}(x)$, получаем

$$\frac{1}{2} R_1 + \frac{1}{2} R_2 \geq C,$$

где R_1 является скоростью передачи канала, отвечающего переходным вероятностям $p_i(j)$, когда выходы имеют вероятности P_i и R_2

является аналогичной величиной для канала, отвечающего переходным вероятностям $p_i(j')$. Эти скорости, конечно, меньше соответственно, чем C_1 или C_2 , так как пропускные способности являются максимальными возможными скоростями. Отсюда получаем желаемый результат

$$\frac{1}{2} C_1 + \frac{1}{2} C_2 \geq C.$$

Различные результаты, полученные нами, могут быть объединены следующим образом.

Т е о р е м а. Конечный дискретный канал без памяти обладает следующими свойствами.

1. Скорость передачи R является строго выпуклой кверху функцией от вероятностей полученных букв Q_i .

2. Скорость R является выпуклой кверху функцией от вероятностей букв P_i на входе.

3. Область пространства вероятностей букв на входе, в которой достигается пропускная способность канала, является выпуклым множеством.

4. Не существует локального максимума R , который не является абсолютным максимумом C .

5. Любая буква на входе, внутренняя для выпуклого тела, определенного другими буквами на входе, может быть выкинута без снижения пропускной способности канала.

6. Для достижения пропускной способности канала нужно использовать только q (соответствующим образом выбранных) букв на входе, где q равно рангу матрицы $\|p_i(j)\|$. Более того, $C \leq \log q$ (по данным Мурога).

7. Пропускная способность является выпуклой вниз функцией переходных вероятностей $p_i(j)$.

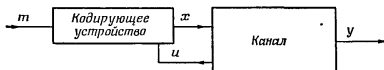
КАНАЛЫ С ДОПОЛНИТЕЛЬНОЙ ИНФОРМАЦИЕЙ НА ПЕРЕДАТЧИКЕ¹⁾

Краткое содержание

В некоторых системах связи, которые передают информацию из одного места в другое, на передающем конце может иметься дополнительная информация. Эта дополнительная информация относится к состоянию передающего канала и может быть использована при кодировании и передаче информации. В этой работе исследуется некоторый тип каналов с дополнительной информацией и определяется их пропускная способность.

Введение

Каналы с обратной связью²⁾, идущей из пункта приема на передающий конец, составляют частный случай каналов, в которых



Р и с. 1.

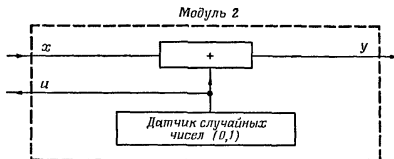
на передающем конце имеется дополнительная информация, которая может быть использована для передачи в прямом направлении. Канал, изображенный на рис. 1, имеет вход x и выход y .

Канал имеет второй выход u , несущий на передающий конец информацию, которая может быть использована в процессе кодирования. Таким образом, кодирующее устройство в качестве входов имеет подлежащее передаче сообщение m и дополнительную информацию u . Последовательность входных букв канала x будет функцией доступной части (т. е. той части, которая относится к прошлому вплоть до текущего момента) этих сигналов.

¹⁾ Shannon C., Channels with side information at the transmitter, *IBM Res. Developm.*, 2, № 4 (1948), 289—293.

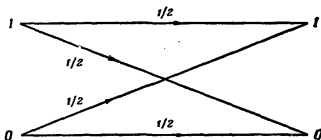
²⁾ Shannon C., Zero error capacity of a noisy channel. (Русский перевод см. стр. 464 данного сборника.—Прим. ред.)

Сигнал u может быть принятым сигналом y , он может быть искаженным шумом — вариантом этого же сигнала, наконец, он может быть не связан с y , а быть лишь статистически коррелированным



Р и с. 2.

с общим состоянием канала. Практическим примером могла бы служить станция, имеющая приемник для измерения текущих шумов на различных частотах. Результаты измерения могли бы быть использованы для выбора частоты, на которой ведется передача.

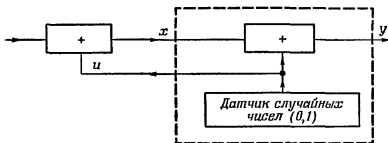


Р и с. 3.

Простой дискретный канал с дополнительной информацией представлен на рис. 2. В этом канале x , y и u являются двоичными переменными; они принимают либо значение нуль, либо значение единица. Канал может быть использован один раз в секунду. Немедленно после каждого использования стохастический прибор выбирает независимо от предыдущих выборов нуль или единицу с вероятностями $1/2$, $1/2$. Эта величина u появляется затем на передающем конце. Следующее передаваемое x суммируется в канале с величиной u по модулю 2, давая на выходе y . Если бы на передатчике не имелось дополнительной информации (соответствующий канал изображен на рис. 3), то канал был бы таким, в котором 0 на входе с вероятностью $1/2$ воспринимался бы на выходе как 0 и с вероятностью $1/2$ как 1; то же самое имело бы место для случая 1 на входе.

Пропускная способность такого канала равна нулю. В случае, однако, когда имеется дополнительная информация, по такому каналу возможно передавать один бит информации в секунду. Информация u используется для того, чтобы компенсировать шум в канале с помощью предварительного обращения нуля и единицы, как это показано на рис. 4.

Не ставя проблему дополнительной информации в ее полной общности, что включало бы случаи, когда приходится учитывать



Р и с. 4.

эффекты, связанные с состоянием канала в предыдущие моменты времени, случаи бесконечного входного и выходного алфавитов и пр., рассмотрим лишь частный случай, для которого будет найдено простое решение.

Дискретный канал без памяти с дополнительной информацией о состоянии

Рассмотрим канал, который имеет конечное число возможных состояний $s_1, s_2, \dots, s_n$. При каждом использовании канала выбирается новое состояние; состояние s_i выбирается с вероятностью g_i . Этот выбор является статистически независимым от предыдущих состояний и предыдущих входных и выходных букв в канале. Состояние известно на передающем конце, оно представляет собой дополнительную информацию. Когда имеет место состояние s_i , канал действует подобно обычному дискретному каналу K_i . Поэтому его действие определяется совокупностью переходных вероятностей $p_{ij}(t)$ ($t = 1, 2, \dots, h$; $i = 1, 2, \dots, a$; $j = 1, 2, \dots, b$), где a — число входных букв, а b — число выходных букв. В совокупности, таким образом, канал описывается множеством вероятностей состояний g_i и переходными вероятностями $p_{ij}(t)$, где g_i — вероятность t -го состояния, а $p_{ij}(t)$ — условная вероятность того, что будет принято j , если при t -м состоянии канала было передано i .

Блочный код для M сообщений (перенумерованных числами $1, 2, \dots, M$) может быть для такого канала дополнительной инфор-

мацией определен следующим образом. Это определение, между прочим, аналогично тому, которое было дано ранее для канала с обратной связью ¹⁾. Пусть n — длина блокового кода. Тогда существуют n функций $f_1(m; u_1)$, $f_2(m; u_1, u_2)$, $f_3(m; u_1, u_2, u_3)$, ..., $f_n(m; u_1, u_2, \dots, u_n)$. В этих функциях m пробегает множество возможных сообщений и поэтому $m = 1, 2, \dots, M$. Все u_i пробегают все значения алфавита дополнительной информации. В рассматриваемом здесь частном случае каждое из u_i может принимать значения от 1 до g . Каждая функция f_i принимает значения в алфавите входных букв канала x . Величина $f_i(m; u_1, u_2, \dots, u_i)$ является входом x_i , который следует использовать в коде, если сообщением является m , а дополнительная информация вплоть до момента, соответствующего i , состоит из $u_1, u_2, \dots, u_i$. Это является математическим эквивалентом утверждения о том, что код представляет собой правило определения следующей передаваемой буквы по каждому сообщению m и заданной дополнительной информации о состоянии канала от начала блока до настоящего момента. Существенным здесь является то, что при выборе следующей передаваемой буквы x_i могут быть использованы только имеющиеся к данному моменту i (а именно, $m; u_1, u_2, \dots, u_i$), а не дополнительная информация $u_{i+1}, \dots, u_n$, которая еще не получена.

Система декодирования для такого кода представляет собой отображение, или функцию $h(y_1, y_2, \dots, y_n)$ принятых блоков длины n , значением которых являются сообщения m ; поэтому h принимает значения от 1 до M . Таков способ выбора декодированного сообщения, которое ставится в соответствие всему принятому блоку $y_1, y_2, \dots, y_n$.

Для некоторой заданной совокупности вероятностей и для некоторых заданных каналов, кодирующей и декодирующей систем существует поддающаяся вычислению вероятность ошибки P_e , которая является вероятностью того, что некоторое сообщение будет закодировано и принято таким образом, что функция h даст на выходе сообщение, отличное от переданного. В частности, будем интересоваться случаями, когда сообщения равновероятны; каждое имеет вероятность, равную $1/M$. Скоростью передачи информации для такого кода является величина $\log M/n$. Нас интересует пропускная способность канала C , т. е. наибольшая скорость R , такая, что возможно выбрать коды со скоростями, произвольно близкими к скорости R , и с произвольно малой вероятностью ошибки P_e .

Можно отметить, что если бы дополнительная информация о состоянии канала не поступала бы на передающий конец, то канал

¹⁾ Shannon C., Zero error capacity of a noisy channel. (Русский перевод см. стр. 464 данного сборника. — Прим. ред.)

действовал бы подобно каналу без памяти с переходными вероятностями

$$p_i(j) = \sum_i g_i p_{ii}(j).$$

При этом условии, следовательно, пропускная способность могла бы быть вычислена обычными методами, существующими для каналов без памяти. С другой стороны, если бы дополнительная информация имелась как на передающем, так и на приемном концах, то, как легко показать, пропускная способность в этом случае выражалась бы формулой $C_2 = \sum_i g_i C_i$, где C_i — пропускная способность канала без памяти с переходными вероятностями $p_{ii}(j)$. Рассматриваемый здесь случай является промежуточным: дополнительная информация имеется на передающем, но не имеется на приемном конце.

Теорема. Пропускная способность дискретного канала K без памяти с дополнительной информацией, которая определяется значениями g_i и $p_{ii}(j)$, равна пропускной способности канала K без памяти (и без дополнительной информации) с тем же самым выходным и входным алфавитом, состоящим из a^h входных букв $X = (x_1, x_2, \dots, x_n)$, где каждое $x_i = 1, 2, \dots, a$. Переходные вероятности $r_x(y)$ для канала K' имеют вид

$$r_x(y) = r_{x_1, x_2, \dots, x_h}(y) = \sum_i g_i p_{ix_i}(y).$$

Любой код и система декодирования для K' могут быть переведены в некоторые эквивалентные код и систему декодирования для K с той же самой вероятностью ошибки. Любой код для K имеет неопределенность сообщений (условную энтропию на букву сообщения принятой последовательности) не большую, чем $R - C$, где C — пропускная способность K' . Любой код со скоростью $R > C$ обладает вероятностью ошибки, ограниченной снизу неравенством (при произвольной длине блока n)

$$P_e \geq \frac{R - C}{6 \left(R + \frac{1}{n} \ln \frac{R}{R - C} \right)}.$$

Можно заметить, что эта теорема сводит исследование данного канала K с дополнительной информацией к исследованию канала K' без памяти, обладающего большим числом входных букв, но не имеющего дополнительной информации. Для определения пропускной способности этого вновь полученного канала K' можно использовать уже известные методы и в результате вычислить пропускную способность исходного канала. Более того, коды для вновь полученного канала могут быть переведены в коды исходного

канала без изменения вероятности ошибки. (На самом деле все статистические свойства кодов остаются теми же самыми.)

Покажем сначала, каким образом коды для канала K' могут быть переведены в коды для канала K . Кодовое слово для вновь полученного канала K состоит из последовательности n букв, принадлежащих входному алфавиту X канала K' . Любая заданная буква X для этого канала может быть представлена как некоторая функция, определенная на алфавите состояний канала и принимающая значения во входном алфавите x канала K . Весь возможный алфавит X состоит из полного множества a^h всех возможных функций, определенных на алфавите состояний с h элементами, со значениями во входном алфавите с h элементами. Таким образом, каждая буква $X = (x_1, x_2, \dots, x_h)$ кодового слова для канала K' может быть интерпретирована как функция от состояния u , принимающая значения во входном алфавите x . Перевод кодов заключается просто в формировании входа x , даваемого этой функцией от состояния канала: Таким образом, если аргумент — состояние u — принимает значение 1, то по каналу K передается x_1 , если же состояние k , то передается x_k . Другими словами, перевод является простым побуквенным переводом без памяти, привносящей зависимость от предыдущих состояний.

Коды для канала K' представляют собой просто иной способ описания некоторых из кодов для канала K ; а именно таких кодов, в которых следующая входная буква x является функцией лишь сообщения m и текущего состояния u и не зависит от предыдущих состояний.

Можно указать также, что не трудно сконструировать простой физический прибор, который, будучи помещенным перед каналом K , сделал бы его сходным с каналом K' . Этот прибор имел бы алфавит X на одном входе и алфавит состояний на другом (этот вход был бы присоединен к линии u на рис. 1). Его выход принимал бы значения в алфавите x и был бы соединен с линией x на рис. 1. Работа прибора заключалась бы в том, чтобы давать выход x , соответствующий значению функции X от состояния u . Ясно, что если произвести перевод кодов, то статистические характеристики для каналов K и K' окажутся идентичными. Вероятность того, что некоторое входное слово для канала K' будет принято как некоторое фиксированное выходное слово, является той же самой для соответствующей операции в канале K . Это показывает справедливость первой половины теоремы.

Для того чтобы доказать вторую половину теоремы, покажем, что в первоначальном канале K приращение условной энтропии (неопределенности) сообщения m при приеме на приемном конце одной буквы не может превзойти C (пропускную способность нового канала K'). Пусть на рис. 1 m будет сообщение; x , y и u

будут следующей входной буквой, выходной буквой и буквой состояния соответственно. Пусть U будет последовательностью предыдущих состояний u , имевших место от начала кодового блока до настоящего момента (точно до u), а Y будет последовательностью предыдущих выходных букв вплоть до текущей буквы y . Предположим теперь, что задан блочный код для кодирования сообщений. Эти сообщения выбираются из некоторого множества с определенными вероятностями (не обязательно равными). При заданных статистике источника сообщений, системе кодирования и статистике канала эти величины m , x , y , U и Y принадлежат все некоторому вероятностному пространству, и различные вероятности, включенные в последующие вычисления, имеют смысл. Таким образом, неопределенность сообщения $H(m/Y)$ при уже принятом Y представляется в виде

$$H(m/Y) = - \sum_{m, Y} P(m, Y) \log P(m/Y) = -E[\log P(m/Y)].$$

(Символ $E(G)$ здесь и в дальнейшем означает математическое ожидание или среднее значение G по вероятностному пространству.) Приращение неопределенности при приеме следующей буквы y равно

$$\begin{aligned} H(m/Y) - H(m/Y, y) &= -E(\log P(m/Y)) + E(\log P(m/Y, y)) = \\ &= E\left(\log \frac{P(m/Y, y)}{P(m/Y)}\right) = E\left(\log \frac{P(m/Y, y) P(Y)}{P(Y, y) P(m, Y)}\right) = \\ &= E\left(\log \frac{P(y/m, Y)}{P(y)}\right) = -E\left(\log \frac{P(Y, y)}{P(Y) P(y)}\right), \\ H(m/Y) - H(m/Y, y) &\leq E\left(\log \frac{P(y/m, Y)}{P(y)}\right). \end{aligned} \quad (1)$$

Последнее неравенство является справедливым в силу того, что член $E\left(\log \frac{P(Y, y)}{P(Y) P(y)}\right)$ представляет собой среднюю взаимную информацию и, следовательно, является неотрицательным. Заметим теперь, что из требований независимости в нашей первоначальной системе следует, что

$$P(y/x) = P(y/x, m, u, U) = P(y/x, m, u, U, Y).$$

Далее, так как x является однозначной функцией m, u , а U системой кодирования, то его можно опустить при указании переменных в условии

$$\begin{aligned} P(y/m, u, U) &= P(y/m, u, U, Y), \\ \frac{P(y, m, u, U)}{P(m, u, U)} &= \frac{P(y, m, u, U, Y)}{P(m, u, U, Y)}. \end{aligned}$$

В силу того что новое состояние u не зависит от предыдущего, $P(m, u, U) = P(u) P(m, U)$ и $P(m, u, U, Y) = P(u) P(m, U, Y)$.

После подстановки и упрощения получим

$$P(y, u/m, U) = P(y, u/m, U, Y).$$

Суммирование по u дает

$$P(y/m, U) = P(y/m, U, Y).$$

Поэтому

$$H(y/m, U) = H(y/m, U, Y) \leq H(y/m, Y) - \\ - E(\log P(y/m, U)) \leq -E(\log P(y/m, Y)).$$

Используя это в неравенстве (1), получаем

$$H(m/Y) - H(m/Y, y) \leq E\left(\log \frac{P(y/m, U)}{P(y)}\right). \quad (2)$$

Покажем теперь, что $P(y/m, U) = P(y/X)$. Здесь X — случайная величина, описывающая функцию, переводящую u в x , задаваемая операцией кодирования следующего входного символа x в канале. Или, что то же самое, X соответствует входной букве в новом канале K' . Имеем $P(y/x, u) = P(y/x, u, m, U)$. Более того, из задания системы кодирования вытекает наличие функционального соотношения, определяющего следующую входную букву x по заданным m , U и u . Поэтому $x = f(m, U, u)$. Если для некоторых двух заданных пар (m, U) и (m', U') при всех u имеет место равенство $f(m, U, u) = f(m', U', u)$, то отсюда следует, что $P(y/m, U, u) = P(y/m', U', u)$ для всех u и y , так как m , U и u приводят к тем же самым буквам x , что и m' , U' и u . Пользуясь этим, получим условие $P(y/m, U) = \sum_u P(u) P(y/m, U, u) = \sum_u P(u) P(y/m', U', u) = P(y/m', U')$. Иначе говоря, пары (m, U) , приводящие к одной и той же функции $f(m, U, u)$, приводят к одной и той же величине $P(y/m, U)$ или, что эквивалентно, $P(y/m, U) = P(y/X)$.

Возвращаясь теперь к неравенству (2), получаем

$$H(m/Y) - H(m/Y, y) \leq E\left(\log \frac{P(y/X)}{P(y)}\right) \leq \max_{P(X)} E\left(\log \frac{P(y/X)}{P(y)}\right), \\ H(m/Y) - H(m/Y, y) \leq C.$$

Это неравенство и требовалось получить для выражения неопределенности. Для каждой принимаемой буквы неопределенность не может уменьшиться более чем на C — пропускную способность нового канала K' . В частности, в блоковом коде с равновероятными сообщениями $R = 1/n \log M$. Если $R > C$, то при окончании блока неопределенность должна еще оставаться по крайней мере равной $nR - nC$, так как вначале она равна nR и может уменьшаться самое большее лишь на величину C для каждой из n букв.

Если, как показано в приложении, неопределенность на одну букву равна по крайней мере $R - C$, то вероятность ошибки при декодировании ограничена следующим неравенством:

$$P_e \geq \frac{R - C}{6 \left(R + \frac{1}{n} \ln \frac{R}{R - C} \right)}.$$

Таким образом, в случае, когда делается попытка использовать код со скоростью $R > C$, то вероятность ошибки оказывается отличной от нуля, какова бы ни была длина блока n . Этим заканчивается доказательство теоремы.

В качестве примера рассмотрим канал с двумя выходными буквами, произвольным числом a входных букв и произвольным числом h состояний. При этом новый канал K' обладает двумя выходными буквами и a^h входными буквами. Однако в канале, у которого имеются ровно две выходные буквы, для того чтобы достичь пропускной способности, нужно, как показано в работе автора¹⁾, использовать лишь две входные буквы; именно в канале K' использовать только те две буквы, одна из которых обладает максимальной, а другая минимальной вероятностями перехода к одной из выходных букв. Эта пара может быть найдена следующим образом. Переходные вероятности для какой-либо буквы из алфавита канала K' являются средними соответствующих переходных вероятностей для множества букв канала K , взятых по одной для каждого состояния. Ясно, что для того, чтобы максимизировать переходную вероятность к одной из выходных букв, требуется выбрать для каждого состояния букву с максимальной вероятностью перехода в эту выходную букву. Подобно этому, для минимизации нужно выбрать для каждого состояния букву с минимальной вероятностью перехода в указанную букву. При использовании лишь этих двух результирующих букв в канале K' получим, что соответствующий канал даст желаемую пропускную способность. Формально, в случае, когда данный канал в состоянии i имеет вероятность $p_{i1}(1)$ перехода входной буквы i в выходную 1 и вероятность $p_{i1}(2) = 1 - p_{i1}(1)$ перехода к другой выходной букве 2, имеем

$$p_1 = \sum_i g_i \max_i p_{i1}(1),$$

$$p_2 = \sum_i g_i \min_i p_{i1}(1).$$

Канал K' , у которого две входные буквы обладают переходными вероятностями p_1 и $1 - p_1$ и p_2 , $1 - p_2$ соответственно для двух

¹⁾ Shannon C., Geometrische Deutung einiger Ergebnisse bei der Berechnung der Kanalkapazität. (Русский перевод см. стр. 488 данного сборника.— Прим. ред.)

выходных букв, имеет пропускную способность, совпадающую с пропускной способностью первоначального канала K .

Рассмотрим еще следующий пример канала с тремя выходными буквами, двумя входными буквами и тремя состояниями. Предположим, что вероятность каждого из состояний равна $1/3$ и матрицы вероятностей для этих трех состояний имеют вид

Состояние 1	Состояние 2	Состояние 3
1 0 0	0 1 0	0 0 1
0 $1/2$ $1/2$	$1/2$ 0 $1/2$	$1/2$ $1/2$ 0

В этом случае в новом канале K' имеются $2^3 = 8$ входных букв. Они образуют следующую матрицу:

$1/2$	$1/2$	0
0	$1/2$	$1/2$
$1/2$	0	$1/2$
$2/3$	$1/6$	$1/6$
$1/6$	$2/3$	$1/6$
$1/6$	$1/6$	$2/3$
$1/3$	$1/3$	$1/3$
$1/3$	$1/3$	$1/3$

Если имеются лишь три выходные буквы, то нужно использовать только три входные буквы, чтобы достичь пропускной способности канала. В рассматриваемом случае, как легко показать, можно (а фактически и необходимо) использовать первые три буквы. Из-за симметрии эти три буквы должны быть употреблены с равными вероятностями; результирующая пропускная способность канала в этом случае равна $\log (3/2)$.

Нетрудно заметить, что если бы в исходном канале не имелось бы информации о состоянии, то канал работал бы подобно каналу с матрицей вероятностей перехода

$1/3$	$1/3$	$1/3$
$1/3$	$1/3$	$1/3$

Ясно, что такой канал обладает нулевой пропускной способностью. С другой стороны, если бы на приемном конце или *и* на приемном, *и* на передающем концах имелась бы информация о предыдущем состоянии, то обе входные буквы могли бы быть полностью различимы и пропускная способность равнялась бы $\log 2$.

П р и л о ж е н и е

Л е м м а. Пусть имеются M возможных событий с вероятностями p_i ($i = 1, 2, \dots, M$). Допустим, что энтропия H удовлетворяет неравенству

$$H = - \sum p_i \ln p_i \geq \Delta;$$

тогда общая вероятность P_e всех возможностей, за исключением наиболее вероятной, удовлетворяет неравенству

$$P_e \geq \frac{\Delta}{6 \ln \left(\frac{M \ln M}{\Delta} \right)}.$$

Доказательство. При заданной энтропии H минимум P_e достигается в случае, когда все вероятности, кроме наибольшей, равны между собой. Это следует из свойства выпуклости энтропии; уравнивание двух вероятностей увеличивает энтропию. Следовательно, можно предположить в наихудшем случае, что имеются $M-1$ возможностей, каждая с вероятностью q , и одна возможность с вероятностью $1 - (M-1)q$. Наше заданное условие перейдет тогда в условие вида

$$-(M-1)q \ln q - [1 - (M-1)q] \ln [1 - (M-1)q] \geq \Delta.$$

Из того, что функция $f(x) = -(1-x) \ln(1-x)$ является выпуклой вниз с наклоном 1 при $x=0$, $f'(x) = 1 + \ln(1-x)$; $f''(x) = -\frac{1}{1-x} \leq 0$ для $0 \leq x \leq 1$, можно заключить, что $f(x) \leq x$ и второе слагаемое в написанном выше неравенстве ограничено величиной $(M-1)q$. Из заданного условия следует поэтому, что

$$-(M-1)q \ln q + (M-1)q \geq \Delta$$

или

$$(M-1)q \ln \frac{e}{q} \geq \Delta.$$

Предположим теперь, в противоположность утверждению леммы, что

$$P_e = (M-1)q < \frac{\Delta}{6 \left(\ln M + \ln \frac{\ln M}{\Delta} \right)}.$$

Так как $-q \ln \frac{e}{q}$ является монотонно возрастающей функцией от q , то отсюда следует, что

$$\begin{aligned} (M-1)q \ln \frac{e}{q} &< \frac{\Delta}{6 \left(\ln M + \ln \frac{\ln M}{\Delta} \right)} \ln \frac{6e(M-1) \left(\ln M + \ln \frac{\ln M}{\Delta} \right)}{\Delta} = \\ &= \frac{\Delta}{6} \left[\frac{\ln \frac{M-1}{\Delta}}{\ln M + \ln \frac{\ln M}{\Delta}} + \frac{\ln 6e}{\ln M + \ln \frac{\ln M}{\Delta}} + \frac{\ln \left(\ln M + \ln \frac{\ln M}{\Delta} \right)}{\ln M + \ln \frac{\ln M}{\Delta}} \right] \leq \\ &\leq \frac{\Delta}{6} \left[1 + 3 + \frac{1}{e} \right] < \Delta \quad (M > 1). \end{aligned}$$

Первая оценка сверху получается с помощью представления соответствующего слагаемого в виде $(\ln \ln M - \ln \Delta + \ln(M-1)) -$

$-\ln \ln M)/(\ln \ln M - \ln \Delta + \ln M)$. Так как $\ln M \geq \Delta$, то нетрудно заметить, что при $M \geq 2$ эта величина ограничена сверху числом 1 (когда $M = 1$, справедливость леммы тривиальна в силу того, что при этом $\Delta = 0$). Оценка соответствующего слагаемого числом 3 очевидна. Последнее слагаемое имеет вид $\ln z/z$. Дифференцируя, устанавливаем, что при $z = e$ она принимает максимальное значение, равное $1/e$. Так как наш результат противоречит условию леммы, то тем самым она доказана.

Основным применением этой леммы является установление нижней грани для вероятности ошибки в системах кодирования. Если известно, что в некоторой ситуации «неопределенность», т. е. условная энтропия сообщения при заданном принятом сигнале, превышает величину Δ , то доказанная лемма дает нижнюю грань для вероятности ошибки. В самом деле, неопределенность является некоторым средним, взятым по множеству принятых сигналов. Отсюда $\Delta = \sum P_i \Delta_i$, где P_i — вероятность принимаемого сигнала i , а Δ_i — соответствующая энтропия сообщения. Если $f(\Delta)$ — это нижняя грань, используемая в лемме, т. е.

$$f(\Delta) = \frac{\Delta}{6 \ln \left(\frac{M \ln M}{\Delta} \right)},$$

то нижней гранью P_e будет $P_e \geq \sum P_i f(\Delta_i)$. Далее, функция $f(\Delta)$ является выпуклой вниз (ее вторая производная не отрицательна в области возможных значений). Отсюда $\sum P_i f(\Delta_i) \geq f(\sum P_i \Delta_i) = f(\Delta)$, и можно заключить, что граница, установленная леммой, остается справедливой даже в более общем случае после простой подстановки среднего значения Δ .

Одним из распространенных случаев, в которых возможно использовать этот результат, является передача кодом со скоростью R , превышающей пропускную способность канала. Во многих случаях это приводит к неопределенности $\Delta = n(R - C)$ после передачи n букв. В этих случаях после подстановки указанных величин в неравенство леммы можно сказать, что для переданного блока вероятность ошибки следующим образом ограничена снизу:

$$P_e \geq \frac{R - C}{6 \left(R + \frac{1}{n} \ln \frac{R}{R - C} \right)} = \frac{R - C}{6 \left(R - \ln \left(1 - \frac{C}{R} \right) \right)}.$$

При указанных условиях это выражение является нижней границей для вероятностей ошибок при скоростях, превышающих пропускную способность канала.

НЕКОТОРЫЕ РЕЗУЛЬТАТЫ ТЕОРИИ КОДИРОВАНИЯ ДЛЯ КАНАЛОВ С ШУМАМИ¹⁾

Краткое содержание

В работе развиваются и уточняются некоторые положения теории кодирования для каналов связи с шумами. Прежде всего с помощью уточнения рассуждений, основанных на методе «случайного» кодирования, дается оценка сверху вероятности ошибки при оптимальном кодировании в случае конечного дискретного канала без памяти. Затем выводится уравнение, позволяющее определить пропускную способность канала с конечным числом состояний в том случае, когда состояния могут быть вычислены и на передающем и на приемном концах. Анализируется также более сложный случай, когда состояние вычислимо на передающем конце, но не обязательно вычислимо на приемном конце.

Оценка вероятности ошибки для дискретного конечного канала без памяти

Дискретный конечный канал без памяти, у которого входной и выходной алфавиты состоят из конечного числа символов (букв), задается совокупностью переходных вероятностей

$$p_i(j), \quad i = 1, 2, \dots, a; \quad j = 1, 2, \dots, b,$$

таких, что $\sum_j p_i(j) = 1$ и все $p_i(j) \geq 0$. Здесь $p_i(j)$ означает вероятность того, что на выходе будет принята буква j , если на входе передавалась буква i . Кодовое слово длины n определяется как последовательность n букв на входе (т. е. как n целых чисел, каждое из которых выбрано из совокупности чисел $1, 2, \dots, a$). Блочный код из блоков длины n с M возможными словами — это отображение совокупности целых чисел от 1 до M (номеров последовательных сигналов) в множество кодовых слов длины n . Система декодирования для такого кода — это отображение всех

¹⁾ Shannon C. E., Certain results in coding theory for noisy channels, *Information and Control*, 1 (1957), 6—25.

выходных слов длины n в совокупность целых чисел $1, \dots, M$, т. е. это некоторая процедура, позволяющая решить, каков был первоначальный номер переданного сообщения.

Будем предполагать, что все целые числа от 1 до M используются с одной и той же вероятностью $1/M$. Для некоторого кода и системы декодирования определим вероятность ошибки P_e как вероятность того, что переданное число будет принято как слово, которому соответствует другое число, т. е. как вероятность декодирования переданного сообщения как другое сообщение. Таким образом,

$$P_e = \sum_u \sum_{v \in S_u} \frac{1}{M} Pr(v|u),$$

где u пробегает все целые числа от 1 до M , v пробегает все принимаемые слова длины n , а S_u — совокупность таких слов, декодируемых не как u . Символ $Pr(v|u)$ означает вероятность получить на выходе v , если было передано сообщение u . Так, если u соответствует слову $(i_1, i_2, \dots, i_n)$ на входе, а v — слово $(j_1, j_2, \dots, j_n)$ на выходе, то

$$Pr(v|u) = p_{i_1}(j_1) p_{i_2}(j_2) \dots p_{i_n}(j_n).$$

В то время как при кодировании мы предположили, что все сообщения имеют равные вероятности $1/M$, при изучении каналов полезно бывает приписать разным словам на входе разные вероятности. В самом деле предположим, что в данном канале мы произвольно задали вероятности различных слов u длины n на входе и что $P(u)$ — вероятность, приписанная слову u . Тогда для вероятностей всех пар, состоящих из слова длины n на входе и слова длины n на выходе, будем иметь

$$Pr(u, v) = P(u) Pr(v|u),$$

где u и v — слова длины n на входе и выходе, а $Pr(v|u)$ — вероятность того, что на выходе будет принято слово v , если было передано слово u (т. е. $Pr(v|u)$ — это произведение переходных вероятностей для соответствующих букв, входящих в состав слов u и v). Если вероятности $P(u)$ заданы, то любая числовая функция от u и v становится случайной величиной. В частности, случайной величиной становится $I(u, v)$ — удельная (отнесенная к одной букве) взаимная информация u и v :

$$I(u, v) = \frac{1}{n} \log \frac{Pr(u, v)}{P(u)Pr(v)} = \frac{1}{n} \log \sum_u \frac{Pr(v|u)}{P(u)Pr(v)}.$$

Обозначим через $q(x)$ функцию распределения этой случайной величины, т. е. положим

$$q(x) = \Pr [I(u, v) \leq x].$$

Функция $q(x)$, конечно, зависит от сделанного нами выбора вероятностей $P(u)$. Сейчас мы докажем теорему, позволяющую при помощи функции $q(x)$ оценить вероятность ошибки для некоторого возможного кода.

Теорема 1. Пусть некоторые вероятности $P(u)$ для слов и длины n на входе приводят к функции распределения $q(I)$ информации на букву. Тогда для любого заданного целого числа M и любого $\theta > 0$ существуют такой блочный код с M возможными словами и такая система декодирования, что если возможные сообщения передаются с равными вероятностями, то вероятность ошибки ограничена сверху так:

$$P_e \leq q(R + \theta) + e^{-n\theta},$$

где $R = (1/n) \log M$.

Доказательство. Для данных M и θ рассмотрим пары (u, v) слов на входе и выходе и определим множество T , состоящее из тех пар, для которых $\log [Pr(u, v)/P(u) Pr(v)] > n(R + \theta)$. Если u выбраны с вероятностями $P(u)$, то вероятность того, что пара (u, v) принадлежит множеству T , равна $1 - q(R + \theta)$.

Рассмотрим теперь ансамбль кодов, образованных следующим образом. Числам $1, 2, 3, \dots, M = e^{nR}$ сопоставляются независимо друг от друга различные возможные слова на входе $u_1, u_2, \dots, u_B$ с вероятностями, соответственно равными $P(u_1), P(u_2), \dots, P(u_B)$. Эта операция приводит к ансамблю кодов, каждый из которых использует M (или меньше) слов на входе. Если имеются B различных слов u_i , то это множество содержит в точности B^M различных кодов, отвечающих B^M различным способам сопоставления M числам B слов на входе. Коды имеют различные вероятности. Так, например, код, в котором всем числам сопоставлено одно и то же слово на входе u_1 (крайне вырожденный случай) имеет вероятность $[P(u_1)]^M$. Код, в котором d_k числам соответствует слово u_k , имеет вероятность $\prod_k [P(u_k)]^{d_k}$. Мы будем иметь дело

со средней вероятностью ошибки для этого множества кодов. Под последней мы будем понимать среднюю вероятность ошибки, когда коды взвешены согласно только что определенным вероятностям. Предположим, что при употреблении любого из этих кодов каждое число передается с вероятностью $1/M$. Отметим, что для некоторых выборов несколько чисел могут соответствовать одному и тому же слову на входе. Тогда это слово будет употребляться с большей вероятностью, чем другие.

Для произвольного кода из нашего ансамбля определим процедуру декодирования следующим образом. Любое принятое v декодируется как число, имеющее наибольшую условную вероятность быть принятым как v . Если эта условная вероятность одинакова для нескольких чисел, мы (по условию) декодируем v как меньшее из них. В силу того что все числа имеют одну и ту же безусловную вероятность передачи $1/M$, при декодировании выбирается одно из тех чисел, для которых вероятность того, что они преобразуются в результате передачи в данное v на выходе, наибольшая.

Вычислим теперь среднюю по нашему ансамблю кодов вероятность ошибки или «неопределенности» P_a . При этом будем пессимистически включать в число ошибок все случаи, когда имеется несколько равновероятных возможностей для получения v .

В произвольном фиксированном коде из нашего ансамбля кодов некоторое слово на входе u или пара (u, v) , вообще говоря, будут встречаться с вероятностями, отличными от $P(u)$ и $Pr(u, v)$. Однако в среднем по этому ансамблю каждое слово u имеет вероятность $P(u)$ и каждая пара (u, v) — вероятность $Pr(u, v)$, так как числам сопоставлялись слова u с вероятностями, в точности равными $P(u)$. И действительно, какому-либо фиксированному сообщению, скажем числу 1, сопоставлялось слово u с вероятностью $P(u)$. Пусть теперь в некотором частном случае числу 1 сопоставлялось u , а было принято v . При этом будет иметь место ошибка или неопределенность, если в рассматриваемом коде существует одно или более чисел, отображаемых внутрь множества $S_v(u)$, состоящего из всех таких слов на входе, что для них вероятность приема на выходе слова v больше или равна соответствующей вероятности для u . В силу того что слова на входе сопоставлялись числам независимо друг от друга, легко подсчитать долю кодов, для которых имеет место только что описанное положение. В самом деле, пусть

$$Q_v(u) = \sum_{u' \in S_v(u)} P(u').$$

Таким образом, $Q_v(u)$ — это вероятность, приписанная совокупности всех слов, для которых условная вероятность приема на выходе слова v больше или равна соответствующей вероятности для u . Доля кодов, в которых число 2 не отображается ни в одно из слов множества $S_v(u)$ (в силу независимости сопоставления слов на входе числам), равна $1 - Q_v(u)$. Доля кодов, для которых в $S_v(u)$ нет элементов, соответствующих другим числам, равна $(1 - Q_v(u))^{M-1}$. Подобное рассуждение приложимо и к любому другому числу так же, как к числу 1. Таким образом, в среднем по ансамблю вероятность ошибки или неопределенности, возникающей, когда при передаче сообщения, соответствующего слову

на входе u , оно принимается как v , в точности равна выражению

$$Pr(u, v) [1 - (1 - Q_v(u))^{M-1}].$$

Отсюда для средней вероятности ошибки или неопределенности получается выражение

$$P_a = \sum_{u, v} Pr(u, v) [1 - (1 - Q_v(u))^{M-1}]. \quad (1)$$

Постараемся теперь оценить эту вероятность при помощи функции распределения q величины информации. Разобьем сначала сумму на две части: сумму по определенному выше множеству T пар (u, v) , для которых $\log [Pr(u, v)/P(u) Pr(v)] > n(R + \theta)$, и сумму по дополнительному к T множеству $\bar{T}$:

$$P_a = \sum_T Pr(u, v) [1 - (1 - Q_v(u))^{M-1}] + \\ + \sum_{\bar{T}} Pr(u, v) [1 - (1 - Q_v(u))^{M-1}].$$

Так как выражение $1 - (1 - Q_v(u))^{M-1}$ представляет собой вероятность, то мы только увеличим первую сумму, если заменим это выражение на 1. После этого первый член становится равным $\sum_T Pr(u, v)$, что по определению совпадает с $q(R + \theta)$. Чтобы оценить вторую сумму, воспользуемся прежде всего тем, что в силу хорошо известного неравенства $(1 - Q_v(u))^{M-1} \geq 1 - (M-1)Q_v(u)$ эта сумма только увеличится от замены в ней

$$1 - (1 - Q_v(u))^{M-1}$$

на $(M-1)Q_v(u)$ и станет еще больше после замены последнего выражения на $MQ_v(u)$. Итак,

$$P_e \leq P_a \leq q(R + \theta) + M \sum_T Pr(u, v) Q_v(u).$$

Покажем теперь, что для u и v , принадлежащих T , имеет место неравенство $Q_v(u) \leq e^{-n(R+\theta)}$. В самом деле, для u и v из T

$$\log \frac{Pr(v|u)}{Pr(v)} > n(R + \theta),$$

т. е.

$$Pr(v|u) > Pr(v) e^{n(R+\theta)}.$$

Если $u' \in S_v(u)$, то

$$Pr(v|u') \geq Pr(v|u) > Pr(v) e^{n(R+\theta)},$$

$$Pr(u', v) > Pr(u') Pr(v) e^{n(R+\theta)},$$

$$Pr(u'|v) > Pr(u') e^{n(R+\theta)}.$$

Суммирование обеих сторон последнего неравенства по $u' \in S_v(u)$ дает

$$1 \geq \sum_{u' \in S_v(u)} Pr(u' | v) > e^{n(R+\theta)} Q_v(u).$$

(Первое из этих неравенств вытекает из того, что сумма вероятностей несовместимых событий не может превосходить единицы.) Мы получили тем самым, что

$$Q_v(u) < e^{-n(R+\theta)}, \quad (u, v) \in T.$$

Используя это неравенство в нашей оценке величины P_e , найдем, что

$$P_e < q(R+\theta) + e^{nR} e^{-n(R+\theta)} \sum_T Pr(u, v) \leq q(R+\theta) + e^{-n\theta};$$

при этом опять учтено то обстоятельство, что сумма вероятностей несовместимых событий не превосходит единицы. Так как средняя по ансамблю кодов вероятность P_e удовлетворяет неравенству $P_e \leq q(R+\theta) + e^{-n\theta}$, то должен существовать некоторый частный код, удовлетворяющий этому же неравенству. Это завершает доказательство.

Теорема 1 является одним из тех результатов, которые указывают на тесную связь между вероятностью ошибки для кодов в каналах с шумом и функцией распределения $q(x)$ величины взаимной информации. Теорема 1 показывает, что если, выбрав вероятность $P(u)$ для слов на входе, мы получили некоторую функцию распределения $q(x)$, то можно создать код, вероятность ошибки для которого оценивается при помощи $q(x)$. Сейчас установим некоторую обратную зависимость: свяжем с $q(x)$ произвольный заданный код. При этом покажем, что вероятность ошибки для кода (при оптимальном декодировании) тесно связана с соответствующей функцией $q(x)$.

Теорема 2. Пусть некоторый фиксированный код имеет $M = e^{nR}$ равновероятных сообщений и пусть функция распределения удельной взаимной информации I между сообщениями и принятыми словами равна $q(x)$. Тогда оптимальная система приема для этого кода приводит к вероятности ошибки P_e , удовлетворяющей неравенствам

$$\frac{1}{2} q\left(P - \frac{1}{n} \log 2\right) \leq P_e \leq q\left(R - \frac{1}{n} \log 2\right).$$

Отметим, что q имеет здесь несколько иной смысл, чем в теореме 1. Здесь она связана с взаимной информацией между сообщениями и принятыми словами, а в теореме 1 — с взаимной информацией между словами на входе и принятыми словами. Если, как это обычно

бывает, все сообщения сопоставляются при кодировании различным словам на входе, то в обоих случаях приходим к одной и той же величине.

Доказательство. Установим сначала оценку снизу. Из определения функции q следует, что вероятность выполнения неравенства

$$\frac{1}{n} \log \frac{Pr(u, v)}{Pr(u) Pr(v)} \leq R - \frac{1}{n} \log 2$$

равна $q [R - (1/n) \log 2]$. Здесь u — сообщение, а v — принятое слово. Но наше неравенство равносильно неравенству

$$Pr(u|v) \leq Pr(u) e^{nR} \frac{1}{2},$$

или (так как $Pr(u) = e^{-nR}$) неравенству

$$Pr(u|v) \leq \frac{1}{2}.$$

Рассмотрим теперь пары (u, v) , для которых неравенство

$$Pr(u|v) \leq \frac{1}{2}$$

имеет место, и представим себе соответствующие линии, соединяющие точки u и v , окрашенными в черный цвет, а все другие линии, соединяющие u с v , — в красный¹⁾. Разделим теперь точки v на два класса. Класс C_1 состоит из тех v , которые декодируются в такие u , что u и v соединены красной линией (а также из тех v , которые декодируются в u , вообще не соединенные с v). Класс C_2 составляют те v , которые декодируются в u , соединенные с v черной линией. Мы установили, что с вероятностью $q [R - (1/n) \log 2]$ пара (u, v) будет соединена черной линией. Те v , которые включены в эти пары, будут попадать в два класса, C_1 и C_2 , с вероятностями, скажем, q_1 и $q_2 = q [R - (1/n) \log 2] - q_1$. Если при этом v окажется в C_1 , то произойдет ошибка, так как на самом деле u было соединено с v черной линией, а декодирование приводит к u , соединенному с v красной линией (или вовсе не соединенному). Следовательно, рассматриваемые случаи приводят к ошибке, возникающей с вероятностью q_1 . Если же наше v окажется принадлежащим классу C_2 , то $Pr(u|v) \leq 1/2$. Это означает, что по крайней мере с такой же вероятностью эти v могут быть получены не из рассматриваемых здесь,

¹⁾ Здесь имеется в виду графическое представление дискретного канала, при котором символы u на входе и символы v на выходе изображаются в виде двух рядов точек и точки u и v , для которых $P(u|v) > 0$ соединяются линиями. См., например, рис. 9, 10 и 11. — *Прим. ред.*

а из других u . Если сложить для этих v вероятности всех пар $Pr(u, v)$, исключая те пары, которые соответствуют друг другу по системе декодирования, то получим по меньшей мере вероятность $q_2/2$, и все эти случаи будут соответствовать неправильному декодированию. Таким образом, получаем следующее неравенство для вероятности ошибки

$$q_e \geq q_1 + \frac{q_2}{2} \geq \frac{1}{2} q \left(R - \frac{1}{n} \log 2 \right).$$

Установим теперь оценку сверху. Рассмотрим систему декодирования, определенную следующим образом. Если для любого принятого v существует некоторое u , такое, что $Pr(u|v) > 1/2$, то это v декодируется в указанное u . Очевидно, что для данного v не может существовать более чем одно такое u , ибо в противном случае сумма соответствующих событий имела бы вероятность, превосходящую единицу. Если же для данного v не имеется ни одного такого u , то декодирование выберем произвольно. Например, можно условиться, что все такие v декодируются в первое из имеющихся на входе сообщений. Вероятность ошибки при таком декодировании меньше или равна вероятности всех пар (u, v) , для которых $Pr(u|v) \leq 1/2$. Иначе говоря, $P_e \leq \sum_S Pr(u, v)$, где S — совокупность пар (u, v) , для которых $Pr(u|v) \leq 1/2$. Условие $Pr(u|v) \leq 1/2$ эквивалентно условию $Pr(u, v)/Pr(v) \leq 1/2$ или условию

$$\frac{Pr(u, v)}{Pr(u) Pr(v)} \leq \frac{1}{2 Pr(u)} = \frac{1}{2} e^{nR}.$$

Последнее равносильно условию

$$\frac{1}{n} \log \frac{Pr(u, v)}{Pr(u) Pr(v)} \leq R - \frac{1}{n} \log 2.$$

Сумма $\sum_S Pr(u, v)$ по определению является значением функции распределения величины $(1/n) \log [Pr(u, v)/Pr(u) Pr(v)]$ в точке $R - (1/n) \log 2$, т. е.

$$P_e \leq \sum_S Pr(u, v) = q \left(R - \frac{1}{n} \log 2 \right).$$

Оценка вероятности ошибки в терминах производящей функции моментов

Теперь мы получим для оценки, указанной в теореме 1, другое выражение, которое сравнительно легко может быть вычислено по известным параметрам канала. Предположим сначала, что приписанные словам в теореме 1 вероятности $P(u)$ равняются произведению вероятностей букв, составляющих эти слова. Так, если u

состоит из последовательности букв $i_1, i_2, \dots, i_n$, то $P(u)$ есть $P_{i_1} \cdot P_{i_2} \dots, P_{i_n}$. Если v состоит из букв $j_1, j_2, \dots, j_n$, то

$$Pr(v) = Pr(j_1) \cdot Pr(j_2) \dots Pr(j_n),$$

$$Pr(u, v) = Pr(i_1, j_1) \cdot Pr(i_2, j_2) \dots Pr(i_n, j_n).$$

Отсюда

$$I(u, v) = \frac{1}{n} \left[\log \frac{Pr(i_1, j_1)}{Pr(i_1) Pr(j_1)} + \log \frac{Pr(i_2, j_2)}{Pr(i_2) Pr(j_2)} + \dots \right] = \\ = \frac{1}{n} [I_1 + I_2 + \dots + I_n],$$

где через I_k обозначена взаимная информация между k -ми буквами слов u и v .

Различные I в последней формуле являются одинаково распределенными независимыми случайными величинами. Величина $nI(u, v)$ равна сумме n одинаково распределенных независимых случайных величин. Следовательно, возникает ситуация, в которой можно применить центральную предельную теорему. Представляется возможным оценить $q(x)$ с помощью каких-либо неравенств, известных для функции распределения подобной суммы. В частности, можно воспользоваться неравенством, которое получил в 1952 г. Чернов [1] для «хвоста» такого распределения. С помощью простых вычислений, использующих обобщенное неравенство Чебышева, он показал, что распределение таких сумм можно оценить при помощи $\varphi(s)$ — производящей функции моментов, относящейся к какой-либо одной случайной величине. Итак, положим

$$\varphi(s) = E[e^{sI}] = \sum_{i,j} P_i p_i(j) \exp \left[s \log \frac{p_i(j)}{P_k p_k(j)} \right] = \\ = \sum_{i,j} P_i p_i(j) \left[\frac{p_i(j)}{\sum_k P_k p_k(j)} \right]^s.$$

Для наших целей более удобным будет использование логарифма производящей функции $\mu(s) = \log \varphi(s)$ (эту величину иногда называют производящей функцией семиинвариантов). Записанный в наших обозначениях результат Чернова имеет вид

$$q(\mu'(s)) \leq e^{[\mu(s) - s\mu'(s)]n}, \quad s \leq 0.$$

Придавая теперь параметру s какие-либо отрицательные значения, получим отсюда оценку для функции распределения q величины информации, экспоненциально зависящую от n . Легко показать, что если дисперсия исходного распределения положительна, то $\mu'(s)$, а также (при отрицательном s) $\mu(s) - s\mu'(s)$ — коэффициент при n в экспоненциальной функции — являются строго возрастающими монотонными функциями. Действительно, производные этих

выражений, равные соответственно $\mu''(s)$ и $-s\mu''(s)$, существуют; кроме того, используя неравенство Шварца, легко показать, что $\mu''(s)$ положительна.

Теорема 3. Пусть в канале без памяти с конечными входными и выходными алфавитами функция $\mu(s)$ представляет собой производящую функцию семиинвариантов для взаимной информации в случае, когда буквам на входе приписаны некоторые вероятности P_i (для буквы i) и когда переходные вероятности для канала равны $p_i(j)$, т. е.

$$\mu(s) = \log \sum_{i,j} P_i p_i(j) \left[\frac{p_i(j)}{\sum_i P_i p_i(j)} \right]^s.$$

Тогда существуют код и система декодирования длины n , скорость создания информации R и вероятность ошибки P_e , удовлетворяющие неравенствам

$$R \geq \mu(s) - (s-1)\mu'(s),$$

$$P_e \leq 2e^{[\mu(s) - s\mu'(s)]n}, \quad s \leq 0.$$

Если $\mu(s) - (s-1)\mu'(s) \rightarrow R^* > 0$ при $s \rightarrow -\infty$, то для $R \leq R^*$

$$P_e \leq e^{(E^* + R^* - R)n},$$

где $E^* = \lim_{s \rightarrow -\infty} [\mu(s) - s\mu'(s)]$ при $s \rightarrow -\infty$.

Доказательство. Согласно теореме 1, имеем

$$P_e \leq q(R + \theta) + e^{-n\theta} \leq e^{[\mu(s) - s\mu'(s)]n} + e^{-n\theta}, \quad s \leq 0,$$

где s выбрано так, что $\mu'(s) = R + \theta$. Последнее условие может быть выполнено, если θ таково, что соответствующее s отрицательно. Выберем величину θ (которая в остальном произвольна) так, чтобы сделать равными между собой коэффициенты при n в показателях степени. (В силу того что первый член монотонно увеличивается вместе с θ , а второй монотонно уменьшается, легко заметить, что такой выбор θ весьма удобен, так как он минимизирует оценку. Фактически оценка не может быть меньше половины ее значение при выбранном значении θ .) Это условие требует, чтобы выполнялось равенство

$$\mu(s) - s\mu'(s) = -\theta = R - \mu'(s),$$

т. е.

$$R = \mu(s) + (1-s)\mu'(s).$$

Так как показатели степени теперь равны, вероятность ошибки ограничена удвоенным первым членом

$$P_e \leq 2e^{[\mu(s) - s\mu'(s)]n}.$$

Это соотношение справедливо для всех отрицательных s и составляет первый результат теоремы.

Однако в некоторых случаях величина R при $s \rightarrow -\infty$ стремится к положительному пределу. Действительно, $R \rightarrow I_{\min} + \log Pr[I_{\min}]$, а показатель степени в оценке P_e стремится к $\log Pr[I_{\min}]$. Для меньших значений R предельные величины показателей степени не могут быть приравнены друг к другу ни при каком выборе s . Однако можно подобрать θ так, чтобы $R + \theta$ было строго меньше, чем $I_{\min}$, скажем, было равно $I_{\min} - \varepsilon$. В силу того что $q(I_{\min} - \varepsilon) = 0$, вероятность ошибки теперь оценивается неравенством $P_e \leq e^{-n\theta} = e^{-n(I_{\min} - R - \varepsilon)}$. Это справедливо при любом ε , значит, можно составлять коды, для которых оценка имеет место при $\varepsilon = 0$. Итак,

$$P_e \leq e^{-n(I_{\min} - R)}$$

для $R < I_{\min}$. Заметим, что при стремлении R к своему предельному значению в первой оценке (т. е. к величине $I_{\min} + \log Pr[I_{\min}]$) показатели степени в обеих оценках сходятся к одной и той же величине $\log Pr[I_{\min}]$. Однако коэффициент при степенной функции уменьшается: вместо 2 он становится равным 1.

Эти оценки могут быть представлены в другом виде, который, возможно, более нагляден. Определим следующим образом некоторую совокупность «скошенных» вероятностей $Q_s(I)$ для различных величин информации

$$Q_s(I) = \frac{Pr(I) e^{sI}}{\sum_I Pr(I) e^{sI}}.$$

Иначе говоря, первоначальная вероятность величины I увеличивается или уменьшается в e^{sI} раз и образованные величины нормируются так, чтобы дать в сумме единицу. Для больших положительных значений s эта отклоненная совокупность вероятностей $Q_s(I)$ увеличивает вероятности $Pr(I)$ для положительных I и уменьшает их для I отрицательных. Если $s = 0$, то $Q_0(I) = Pr(I)$. При отрицательном s вероятности отрицательных значений I увеличиваются за счет вероятностей положительных величин I . Если $s \rightarrow \infty$, то $Q_s(I) \rightarrow 0$, кроме того случая, когда $I = I_{\max}$; при этом $Q_s(I_{\max}) \rightarrow 1$. Здесь величина $I_{\max}$ — это наибольшее значение I , имеющее положительную вероятность ($I_{\max}$ существует, так как множество пар (u, v) конечно). Эти отклоненные вероятности удобны для оценки «хвостов» распределения, которые являются свертками других

распределений. В терминах $Q_s(I)$ мы можем написать

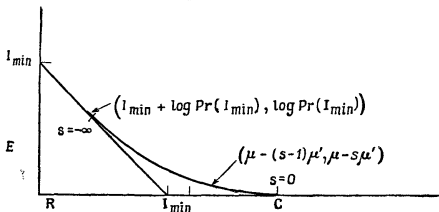
$$\mu(s) = \log \sum_I \text{Pr}(I) e^{sI} = \sum_{I'} Q_s(I') \log \sum_I \text{Pr}(I) e^{sI},$$

$$\mu'(s) = \frac{\sum_I \text{Pr}(I) e^{sI} I}{\sum_I \text{Pr}(I) e^{sI}} = \sum_I Q_s(I) I,$$

$$\mu(s) - s\mu'(s) = \sum_I Q_s(I) \log \frac{\text{Pr}(I)}{Q_s(I)},$$

$$\mu(s) - (s-1)\mu'(s) = \sum_I Q_s(I) \left[I + \log \frac{\text{Pr}(I)}{Q_s(I)} \right].$$

Некоторые интерес представляют значения коэффициентов при n в показателях степени. Они характеризуют скорость приближения P_e к нулю при увеличении n . На рис. 1 изображена их типичная зависимость от R . Мы предполагаем здесь, что вероятности P_i



Р и с. 1.

для букв таковы, что они реализуют пропускную способность канала. Коэффициент E при n для первой оценки в рассматриваемой теореме представляет собой кривую, которая касается оси R в точке C (где $s=0$), кончается, когда $R = I_{\min} + \log \text{Pr}(I_{\min})$, а $E = \log \text{Pr}(I_{\min})$ (т. е. при $s = -\infty$) и обращена своей выпуклостью вниз. Второй оценке в теореме соответствует кривая E , которая представляет собой прямую линию с тангенсом угла наклона, равным -1 . Она проходит через точку $(I_{\min} + \log \text{Pr}(I_{\min}), \log \text{Pr}(I_{\min}))$ и пересекает оси в точках $(I_{\min}, 0)$ и $(0, I_{\min})$. В окрестности точки $R = C$ кривая ведет себя как

$$E \approx \frac{(C-R)^2}{2\mu''(0)},$$

где $\mu''(0)$ — дисперсия I . Все перечисленные свойства вытекают непосредственно из формул для этих кривых.

Предельный показатель степени (при $n \rightarrow \infty$) имеет вид $E = \mu(s) - (s-1)\mu'(s)$. Мы получаем

$$\frac{dE}{dR} = \frac{dE}{ds} / \frac{dR}{ds} = \frac{s}{1-s},$$

откуда видно, что производная функция $E = E(R)$ монотонно уменьшается, изменяясь от 0 до -1 , когда s изменяется от 0 до $-\infty$. Так как второй оценке на графике этой функции соответствует прямая линия с тангенсом угла наклона, равным -1 , то обе оценки не только совпадают по своему значению, но также имеют одну и ту же производную, как показано на рис. 1.

Кривая будет такой, как это указано, если P_i — вероятности, для которых реализуется максимальная скорость передачи, равная пропускной способности канала. При этом

$$R(0) = \mu(0) - (0-1)\mu'(0) = \mu'(0) = C.$$

Однако при использовании соответствующих $\mu(s)$ оценка, выведенная в теореме, справедлива для любой совокупности вероятностей P_i . Для того чтобы получить более сильный результат, оценку нужно сделать оптимальной для каждого значения R с помощью вариации вероятностей P_i . То же самое относится и к прямой линии, где требуется максимизировать $I_{\min}$. Если проделать это, то получится некоторая кривая, которая будет служить огибающей всех возможных кривых этого типа с различными величинами P_i . Так как каждая отдельная кривая выпукла вниз, то то же самое верно для огибающей. Уравнение для огибающей может быть получено решением по методу Лагранжа задачи о максимуме выражения $R + \lambda E + \lambda \sum_i P_i$. Нужно, конечно, иметь в виду,

что P_i должны быть неотрицательными. Эта задача похожа на ту, которая возникает при отыскании пропускной способности канала. Уравнения для огибающей будут иметь вид

$$E = \mu(s) - s\mu'(s),$$

$$R = \mu(s) - (s-1)\mu'(s),$$

$$(1+\lambda) \frac{\partial \mu}{\partial P_i} - (1+\lambda)s \frac{\partial \mu'}{\partial P_i} + \frac{\partial \mu'}{\partial P_i} + \eta = 0$$

для всех i , кроме тех, для которых $P_i = 0$, причем

$$\sum P_i = 1.$$

Полученная оценка должна быть сделана максимальной при помощи выбора различных подмножеств, в которых P_i не обращаются в нуль.

Оценка сверху, полученная в теореме 3, ни в коей мере не является самой сильной среди тех, которые можно найти. Так, когда $n \rightarrow \infty$, коэффициенты при n в показателях степени можно, вообще говоря, улучшить с помощью более тонких рассуждений. Мы надеемся в другой работе развить эти дальнейшие результаты¹⁾, а также получить соответствующую оценку снизу того же самого экспоненциального типа для вероятности ошибки. Несмотря на это, оценка, полученная в теореме 3, имеет два преимущества: она проста и удобна в употреблении. Кроме того, она имеет универсальность, чего недостает некоторым из более сильных результатов (принимаяющих простой вид, лишь когда n очень велико).

Пропускная способность канала с конечным числом состояний, состояние которого вычислимо на обоих концах

В некоторых каналах с памятью внутреннее состояние канала может быть вычислено по начальному состоянию (с которого начинается передача и которое предполагается известным) и последовательности переданных букв. Бывает также, что одновременно можно определить состояние в любое время на приемном конце, если известно начальное состояние и последовательность принятых букв. Условимся говорить, что для таких каналов состояние *вычислимо на обоих концах*.

Чтобы удовлетворить первому требованию, необходимо, очевидно, чтобы для любого (достижимого) внутреннего состояния s следующее состояние t было функцией s и x , т. е. $t = f(s, x)$, где x — переданная буква.

Для того чтобы состояние было вычислимо на приемном конце, необходимо, чтобы для всех достижимых состояний s следующее состояние t было функцией s и принятой буквы y , т. е. $t = g(s, y)$.

Для каждой возможной пары s, t можно найти подмножество $A(s, t)$ значений x , переводящих состояние s в состояние t , и подмножество $B(s, t)$ значений y , соответствующих переходу из состояния s в состояние t . Для каждой буквы на входе x из множества $A(s, t)$ буква y на выходе необходимо будет принадлежать множеству $B(s, t)$ и при этом будут существовать переходные вероятности, т. е. вероятности (в состоянии s) того, что если было передано x , то будет принято y . Можно считать, что для фиксированной пары s, t множества букв $A(s, t)$ и $B(s, t)$ вместе с соответствующими переходными вероятностями определяют некоторый дискрет-

¹⁾ Изложению подобных оценок на основе неопубликованных работ Шеннона посвящена последняя глава книги Фано [Fano R., Transmission of information, MIT Press and Wiley, New York—London, 1961, 384; готовится к печати русский перевод (ИЛ)]. —Прим. ред.

ный канал без памяти, относящийся к паре s, t . Это значит, что рассматривается канал без памяти с входным алфавитом, состоящим из букв множества $A(s, t)$, выходным алфавитом, состоящим из букв множества $B(s, t)$ и соответствующими им переходными вероятностями.

Этот канал может быть физически реализован на основе данного канала следующим образом. Сначала данный канал приводится в состояние s , после чего передается какая-то буква из множества $A(s, t)$ (результатирующее состояние будет t). Затем канал снова возвращается в состояние s и передается вторая буква из множества $A(s, t)$ и т. д. Пропускная способность такого дискретного канала без памяти может быть найдена обычными методами. Обозначим через C_{st} пропускную способность (в натуральных единицах) канала, относящуюся к переходу из состояния s в состояние t . Пусть $N_{st} = e^{C_{st}}$. Таким образом, N_{st} равно числу эквивалентных букв, на которые не влияет шум, для (s, t) -подканала. В случае когда множество $A(s, t)$ пусто, будем считать, что $N_{st} = 0$.

Состояния нашего канала могут быть следующим образом сгруппированы в классы эквивалентности. Состояния s и s' принадлежат одному и тому же классу, если существует последовательность букв на входе, передача которой, начатая в состоянии s , окончательно приводит канал в состояние s' , и обратно: существует некоторая последовательность, переводящая s' в s . Классы эквивалентности можно частично упорядочить следующим образом. Если существует последовательность букв, переводящая элемент одного класса в элемент другого класса, то первый класс считается имеющим более высокий порядок, чем второй класс.

Внутри класса эквивалентности можно рассматривать всевозможные замкнутые последовательности состояний, т. е. всевозможные пути, возвращающие канал, находившийся в начале передачи в некотором определенном состоянии, в то же самое состояние. Число состояний, входящих в такой цикл, будет называться длиной цикла. Наибольший общий делитель всех длин циклов в каком-то фиксированном классе эквивалентности будет называться основным периодом этого класса. Эти структурные свойства аналогичны свойствам цепи Маркова с конечным числом состояний, если вместо «перехода с положительной вероятностью» рассматривать «переход, возможный при некоторой букве на входе».

Будем теперь рассматривать каналы, в которых имеется лишь один класс эквивалентности. Это значит, что имеется возможность перейти из любого состояния s в любое состояние t при помощи некоторой последовательности букв на входе (т. е. любое состояние достижимо из любого другого состояния). Рассмотрение более общего случая нескольких эквивалентных классов только более громоздко, однако оно не приводит к существенным затруднениям.

Теорема 4. Пусть K — канал с конечным числом состояний и конечными алфавитами. Пусть в канале K состояния вычислимы на обоих концах и любое состояние достижимо из любого другого состояния. Обозначим через N_{st} число эквивалентных букв для подканала, связанного с переходами из состояния s в состояние t и пусть N — (единственное) вещественное положительное собственное значение матрицы N_{st} , т. е. действительный положительный корень уравнения

$$|N_{st} - N\delta_{st}| = 0.$$

Тогда N — это эквивалентное число букв для данного канала K , так что его пропускная способность определяется равенством $C = \log N$.

Доказательство. Покажем сначала, что существуют блочные коды, которые позволяют осуществлять передачу с любой скоростью $R < C$ при сколь угодно малой вероятности ошибки. Рассмотрим матрицу N_{st} . Если возвести ее в степень n , то получим новую матрицу, элементы которой обозначим через $N_{st}^{(n)}$. Элемент $N_{st}^{(n)}$ можно рассматривать как сумму произведений; каждое такое произведение соответствует некоторому пути из n отдельных шагов, переводящему состояние s в состояние t , и включает все элементы исходной матрицы, отвечающие этому пути, а сумма произведений распространяется по всем возможным путям. Справедливость этого утверждения легко может быть установлена с помощью метода математической индукции и правила умножения матриц.

Кроме того, $N_{st}^{(n)}$ можно истолковать как эквивалентное число букв для канала без памяти, определенного следующим образом. Представим себе, что начальное состояние исходного канала будет s . Будем употреблять в качестве «букв» на входе последовательности исходных букв длины n , оперируя при этом только теми последовательностями длины n , которые переводят канал в состояние t . «Буквами» на выходе будут последовательности принятых букв длины n , которые могут получиться при этих условиях. Такой канал можно рассматривать как некоторую «сумму» каналов (соответствующих различным последовательностям состояний, начинающимся в s , кончающимся в t и содержащим n шагов), каждый из которых является «произведением» каналов (соответствующих простым переходам из одного состояния в другое). (Суммой двух каналов называется канал, в котором может быть использована в качестве буквы на входе одна из таких букв для любого из двух каналов; произведение двух каналов определяется как канал, в котором в качестве буквы на входе используется упорядоченная пара букв двух первоначальных каналов.) Эквивалентное число букв, свободных от воздействия шума, обладает аддитивным

свойством для суммы каналов и мультипликативным свойством для их произведения. Отсюда следует, что только что описанный канал, который соответствует последовательностям букв длины n , переводящим исходный канал из состояния s в состояние t , имеет эквивалентное число букв, равное элементу матрицы $N_{st}^{(n)}$.

Первоначальная матрица N_{st} является матрицей с неотрицательными элементами. Следовательно, она имеет положительное действительное собственное значение, модуль которого больше или равен модулям остальных собственных значений. Более того, при нашем предположении относительно того, что из любого состояния в любое другое состояние можно перейти с помощью некоторой последовательности букв, существует только одно положительное действительное собственное значение. Если d — наибольший общий делитель длин замкнутых путей (по последовательностям состояний), то должно существовать d собственных значений, равных действительному положительному корню характеристического уравнения, умноженному на корни d -й степени из единицы. Когда матрица возводится в n -ю степень, член $N_{st}^{(n)}$ либо равен нулю (если невозможно перейти из s в t в точности за n шагов), либо асимптотически совпадает с постоянной, умноженной на N^n .

В частности, для n , сравнимого с нулем по модулю d (т. е. делящегося на d), диагональные элементы $N_{ii}^{(n)}$ асимптотически ведут себя как постоянная, умноженная на N^n . В то же время в противном случае эти утверждения являются хорошо известными результатами теории матриц с неотрицательными элементами. Они были впервые получены еще Фробениусом [4] и здесь доказываться не будут.

Если выбрать n достаточно большим кратным d , то $N_{ii}^{(n)} > k N^n$, где k — положительное число. Выбирая n достаточно большим, можно сделать пропускную способность канала, «буквы» на входе которого переводят состояние 1 в то же состояние 1 за n шагов, больше чем $(1/n) \log k N^n = \log N + (1/n) \log k$. В силу того что последний член можно сделать произвольно малым, получаем пропускную способность, как угодно близкую к $\log N$. Так как, разумеется, можно пользоваться исходным каналом таким специальным образом (применяя лишь переходы из состояния 1 в состояние 1 за n шагов), то ясно, что этот исходный канал не может иметь пропускную способность, меньшую чем $\log N$.

Чтобы показать, что эта пропускная способность не может быть превышена, рассмотрим канал K_n , который определяется для последовательностей длины n следующим образом. Когда начинается блок длины n , канал K_n может находиться в произвольном состоянии, выбранном из множества состояний, соответствующ-

щих состояниям канала K . Это достигается выбором «буквы состояния» на передающем конце и передачей ее на приемный конец без искажения ее шумом. Для следующих n символов канал ведет себя, как заданный исходный канал K , т. е. определяется теми же самыми ограничениями и вероятностями. В конце этого блока из n символов начальное состояние для следующего блока выбирается на передающем конце произвольным образом независимо от предыстории канала. Рассматривая блок длины n (включая информацию о начальном состоянии) как одну букву на входе, а соответствующий ему блок y , включающий также принятую первой «букву состояния» как одну букву на выходе, получаем канал без памяти K_n .

Для любой заданной пары s, t начального и конечного состояний соответствующая пропускная способность равна $\log N_{st}^{(n)}$. Так как имеется сумма таких заданных каналов, то пропускная способность K_n равна $\log \sum_{s,t} N_{st}^{(n)}$. Каждое слагаемое в этой

сумме ограничено постоянной, умноженной на N^n . В силу того что имеется только конечное число различных слагаемых (так как имеется лишь конечное число различных состояний), можно выбрать одну и ту же постоянную для всех членов, т. е. положить $N_{st}^{(n)} < kN^n$ (для всех n, s, t). Взяв n достаточно большим, очевидно, получим, что удельная (на одну букву) пропускная способность K_n ограничена величиной $\log N + \epsilon$, где $\epsilon > 0$ можно выбрать сколь угодно малым. Но теперь любой код, который может употребляться в исходном канале, может также быть использован и в канале K_n при любом n , ибо последний имеет те же самые ограничения, кроме тех, которые возникают на концах блоков длины n (в этих точках все ограничения устранены). Отсюда вытекает, что пропускная способность исходного канала меньше или равна пропускной способности канала K_n при всех n и, следовательно, она меньше или равна $\log N$. Это утверждение завершает доказательство теоремы.

Этот результат может быть обобщен в нескольких направлениях. Во-первых, несущественным является предположение о том, что алфавиты конечны. В действительности канал, связанный с переходом из состояния s в состояние t , может быть произвольным каналом без памяти, а не только дискретным каналом с конечным алфавитом.

Второе, не очень существенное обобщение связано с тем, что нет необходимости требовать, чтобы состояние было вычислено на приемном конце после приема каждой буквы, если тем не менее на приемном конце можно определить все предыдущие состояния. Поэтому можно не требовать, чтобы следующее состояние было

некоторой функцией предыдущего состояния и принятой буквы; можно потребовать лишь, чтобы не существовало двух различных последовательностей состояний, переводящих некоторое начальное состояние s в заданное конечное состояние t и совместимых с одной и той же последовательностью принятых букв.

Пропускная способность канала с конечным числом состояний в случае, когда состояние вычислимо на передающем конце и не вычислимо на приемном конце

Рассмотрим теперь канал с конечным входным и выходным алфавитами и с конечным числом внутренних состояний. Пусть, кроме того, имеет место дополнительное условие, согласно которому состояние канала в начале передачи известно на передающем конце и каждое последующее состояние может быть вычислено на этом конце для любой возможной последовательности букв на входе. Иначе говоря, предполагается, что следующее состояние является функцией текущего состояния и текущей буквы на входе. Такой канал определяется переходной функцией для состояний вида $s_{n+1} = f(s_n, x_n)$ (состояние s_{n+1} является функцией состояния s_n и n -го посланного символа) и условными вероятностями $p_{yx}(y)$, заданными для состояния s , определяющими вероятность того, что если была передана буква x , то на выходе будет принята буква y . Мы не предполагаем при этом, что состояние вычислимо и на приемном конце.

Как и раньше, состояния такого канала можно сгруппировать в частично упорядоченные совокупности классов эквивалентности. Будем снова рассматривать каналы, в которых имеется лишь один класс эквивалентности, т. е. будем предполагать, что можно перейти из любого состояния s в любое другое состояние t с помощью некоторой последовательности букв на входе.

Сначала определим пропускную способность для некоторого фиксированного состояния s . Пусть рассматриваемый канал находится в состоянии s и пусть $X_1 = (x_1, x_2, \dots, x_n)$ — последовательность букв на входе, которые в конце концов приводят канал в прежнее состояние s . Если канал находится в состоянии s и используется последовательность X_1 , то можно вычислить условные вероятности различных возможных последовательностей на выходе Y длины n . Так, в случае, когда последовательность X_1 переводит канал через состояния $s, s_2, \dots, s_n, s$, тогда условная вероятность последовательности $Y_1 = (y_1, y_2, \dots, y_n)$ равна $Pr(Y_1/X_1) = P_{sx_1}(y_1) = P_{s_2x_2}(y_2) \dots P_{s_nx_n}(y_n)$. Рассмотрим последовательности X_i (переводящие канал из состояния s в то же состояние s за n

шагов) как отдельные буквы на входе нового канала без памяти. В качестве букв на выходе y этот канал будет иметь последовательности Y , а в качестве переходных вероятностей — соответствующие условные вероятности. Пусть $C(n, s)$ — пропускная способность этого канала, а $C(s)$ — верхняя грань величин $(1/n) C(n, s)$, когда n пробегает всевозможные положительные целые значения. Отметим следующие свойства введенных величин.

1. $C(kn, s) \geq kC(n, s)$. Это следует из того факта, что для достижения наибольшей скорости передачи (т. е. наибольшей пропускной способности канала) выбор вероятностей для цепочек букв на входе X длины kn может использоваться по крайней мере с таким же успехом, как и выбор произведения вероятностей для k последовательностей X , имеющих длину n . Отсюда следует, что если аппроксимировать $C(s)$ с точностью до ε при помощи выбора некоторого фиксированного значения числа n (т. е. если $|C(s) - C(n, s)| < \varepsilon$), то аппроксимация будет не хуже и для бесконечной последовательности $2n, 3n, 4n, \dots$ значений этого числа.

2. $C(s) = C$ не зависит от состояния s . Это доказывается следующим образом. Выберем некоторую последовательность букв на входе U , ведущую из состояния s' в состояние s , и другую последовательность V , ведущую из состояния s в состояние s' . Ни одна из этих последовательностей не содержит более чем m букв, где m — число (конечное) состояний канала. Выберем n_1 таким образом, чтобы выполнялось неравенство $C(n_1, s) > C(s) - \varepsilon/2$ и чтобы в то же время n_1 было достаточно большим, так что

$$\left(C(s) - \frac{\varepsilon}{2}\right) \frac{n_1}{n_1 + 2m} \geq C(s) - \varepsilon.$$

Этого можно добиться, так как, согласно замечанию, приведенному при обсуждении свойства 1, $C(s)$ может быть аппроксимировано сколько угодно точно при использовании произвольно больших значений n_1 .

Сконструируем теперь некоторую совокупность последовательностей X для состояния s' , используя имеющиеся последовательности для состояния s и добавляя к ним вначале последовательность U , а в конце — последовательность V . В случае когда каждой из полученных последовательностей для состояния s' приписана вероятность, равная той, которая приписывалась последовательности X для состояния s с целью достижения пропускной способности $C(n, s)$, скорость передачи для канала из последовательностей s' будет в точности равна $C(n, s)$, но на этот раз длины цепочек будут уже самое большее равны $n_1 + 2m$, т. е. немного больше n_1 . Отсюда следует, что

$$C(s') \geq [C(s) - (\varepsilon/2)] n_1 / (n_1 + 2m) \geq C(s) - \varepsilon.$$

Очевидно, что, поменяв местами s и s' , мы получим обратный результат: $C(s) \geq C(s') - \varepsilon$, и поэтому $C(s) = C(s')$. (Отметим, что если имеется несколько классов эквивалентности, то для каждого из них существует своя пропускная способность C и пропускные способности для различных классов не обязательно совпадают друг с другом.)

3. Пусть $C(n, s, s')$ — пропускная способность, вычисленная для последовательностей длины n , которые начинаются в состоянии s и оканчиваются в s' . Пусть $C(s, s') = \lim_{n \rightarrow \infty} (1/n) C(n, s, s')$. Тогда получается $C(s, s') = C(s) = C$. Это обстоятельство имеет место в силу того, что можно заменить последовательности, переводящие из состояния s в состояние s' , последовательностями, переводящими из s в s при помощи добавления в конце дополнительной последовательности, имеющей длину не больше чем m . Выбирая n достаточно большим, можно сделать эффект добавления последовательности длины m произвольно малым (как и при доказательстве свойства 2), так что $C(s', s') \geq C(s) - \varepsilon$. Точно так же последовательности, переводящие из s в s , с помощью которых аппроксимируют $C(s)$ и которые могут быть сделаны произвольно длинными, переводятся в последовательности, переводящие из s в s' добавлением не более m букв. Поэтому $C(s) \geq C(s, s') - \varepsilon$. Следовательно, $C(s) = C(s, s')$.

Сейчас мы хотим показать, что, начиная с состояния s_1 , можно вести передачу с произвольно малой вероятностью ошибки и со скоростью $R < C$, где C — величина, определенная ранее в формулировке свойства 3. Точнее, нами будет доказано следующее.

Теорема 5. Для любого $R < C$ существует величина $E(R) > 0$, такая, что для любого $n = kd$ (равного некоторому целому числу, умноженному на d — основную длину цикла) найдется блочный код длины n с M возможными словами, для которого $(1/n) \log M \geq R$ и вероятность ошибки удовлетворяют неравенству $P_e \leq e^{-E(R)n}$. Не существует никакой последовательности кодов, для которой с увеличением длины блока вероятность ошибки стремилась бы к нулю и скорость передачи была бы больше C .

Доказательство. Положительное утверждение этой теоремы доказывается следующим образом. Пусть $R_1 = (R + C)/2$. Допустим, что s_1 — начальное состояние канала. Рассмотрим последовательности букв, которые переводят s_1 в s_1 за n_1 шагов. Выберем n_1 так, чтобы $C(n_1, s_1) > (3C + R)/4$. Используем эти последовательности как буквы на входе и построим коды, для которых скорость передачи равна R_1 . Согласно теореме 2, вероятность ошибки будет экспоненциально убывать с ростом длины кода. В нашем случае коды будут иметь длину $n_1, 2n_1, 3n_1, \dots$ исходных букв, но

это лишь уменьшает коэффициент при n в показателе степени в $1/n_1$ раз. Следовательно, положительное утверждение теоремы доказано для чисел, кратных n_1 . Чтобы доказать это же утверждение для всех чисел, кратных d , заметим сначала, что оно справедливо для всех достаточно больших чисел, кратных d . В самом деле, при достаточно больших числах, кратных n_1 , эффект добавления кодовых слов, переводящих состояние снова в s_1 после числа букв, кратного d , может быть сделан произвольно малым (так что скорость передачи фактически не изменится). Для меньших чисел, кратных d , можно использовать теперь любой код с некоторой вероятностью ошибки меньше чем 1 (например, интерпретировать любое полученное слово как сообщение 1 с $P_e = 1 - 1/M < 1$). При этом будем иметь некоторую конечную совокупность кодов вплоть до некоторого числа, кратного d , начиная с которого справедлива единая экспоненциальная оценка. Ясно, что здесь можно будет выбрать некоторый коэффициент $E(R)$, такой, что $P_e < e^{-E(R)n}$ для любого n , кратного d .

Обратное утверждение теоремы (то, что пропускная способность C не может быть превзойдена) доказывается методом, подобным тому, который использовался в случае, когда состояние было вычислимо также и на приемном конце. В самом деле, рассмотрим канал K_n , определенный следующим образом. Вначале данный канал K помещается в произвольное состояние и номер этого состояния передается без помех на приемный конец. Затем передаются n букв при таких же ограничениях и с такими же вероятностями, как у данного канала K . Конечное состояние также передается. Этот процесс повторяется далее блоками длины n . Теперь имеется некоторый канал без памяти, который для любого n «включает в себя» данный канал. В канале K_n можно использовать любое кодирование, предназначенное для канала K , и оно будет приводить к столь же хорошим вероятностям ошибок. Следовательно, для любого n пропускная способность данного канала должна быть меньше или равна пропускной способности канала K_n . С другой стороны, в действительности K_n является «суммой» множества каналов, соответствующих последовательностям, переводящим из состояния s в состояние t за n шагов (т. е. каналов, пропускные способности которых были обозначены ранее через $C(n, s, t)$). Как мы видели выше, для всех достаточно больших n и для всех s и t имеет место неравенство $(1/n) C(n, s, t) < C + \varepsilon$. Следовательно, для всех $n > n_0$ пропускная способность канала K_n меньше, чем $C + \varepsilon + (1/n) \log m^2$, где m — число состояний. Отсюда следует, что пропускная способность канала K не больше, чем C .

Интересно сравнить результаты этого раздела, в котором считалось, что состояния вычислимы только на передающем конце, с результатами, полученными в предыдущем разделе, где пред-

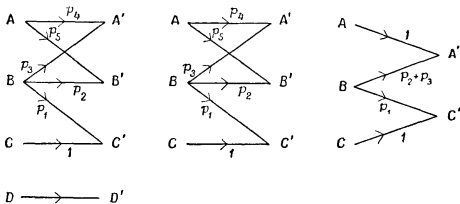
полагалось, что состояния вычислимы на обоих концах. В последнем случае было дано явное выражение для пропускной способности, требующее лишь вычисления пропускной способности каналов без памяти и решения алгебраического уравнения. В первом же случае решение дается в менее явном виде: оно требует отыскания некоторых пределов довольно сложного типа.

ЛИТЕРАТУРА

1. Chernov H., A measure of asymptotic efficiency for tests of a hypothesis based on the sum of observations, *Ann. Math. Statistics*, 23 (1952), 493—507.
2. Элайес П., Кодирование для двух каналов с шумами, в сб. «Теория передачи сообщений», М, ИЛ, 1957, 114.
3. Feinstein A., Error bounds in noisy channels without memory, *IRE Trans. IT-1* (1955), September, 13—14.
4. Frobenius G., Über Matrizen aus nichtnegativen Elementen, *Sitzungsber. Akad. Wiss.*, Berlin (1912), 456—477.
5. Shannon C. E., A mathematical theory of communication. (Русский перевод см. стр. 243 данного сборника.—Прим. ред.)
6. Shannon C. E., The zero error capacity of a noisy channel. (Русский перевод см. стр. 464 данного сборника.—Прим. ред.)

ЗАМЕЧАНИЯ О ЧАСТИЧНОМ УПОРЯДОЧИВАНИИ КАНАЛОВ СВЯЗИ¹⁾

В статье определяется понятие частичного упорядочивания для дискретных каналов без памяти. Это упорядочивание транзитивно и сохраняется при таких операциях над каналами, как сложение и умножение. Основной результат работы состоит в доказательстве



Р и с. 1. Примеры, иллюстрирующие отношение включения.

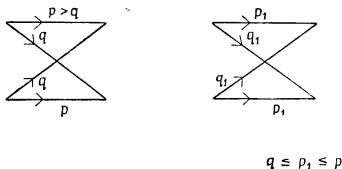
того, что если K_1 и K_2 — два канала и $K_1 \supseteq K_2$, то для любого кода для канала K_2 существует по крайней мере такой же хороший код для канала K_1 , если качество кода измерять вероятностью ошибки.

Рассмотрим три дискретных канала без памяти, изображенные на рис. 1. Про первый можно сказать, что он включает второй, так как первый сводится ко второму, если в нем использовать только буквы A, B, C . Все, что можно передавать по второму каналу, можно с помощью такого искусственного сведения передавать и по первому (но по первому каналу можно, конечно, передать больше, используя полный алфавит). Второй в некотором смысле включает третий, так как третий можно получить из второго, объединив буквы A' и B' на выходе в одну и ту же букву. Можно представить себе как бы некоторое приспособление, поставленное

¹⁾ Shannon C., A note on a partial ordering for communication channels, *Information and Control*, 1, № 4, December (1958), 390.

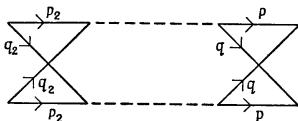
на выходе второго канала, которое выдает букву A' , когда на выходе появляется буква A' или B' и пропускает букву C' без изменения.

Это и есть примеры отношения включения каналов, которое будет определено и рассмотрено ниже. Другой пример представлен



Р и с. 2. Еще один пример включения.

на рис. 2 двумя двоичными симметричными каналами. В этом случае первый канал можно свести ко второму не с помощью метода идентификации букв, а с помощью добавления некоторого статистического устройства на входе или на выходе; а именно, если поместить, как показано на рис. 3, перед первым каналом (или после



Р и с. 3. Сведение левого канала, изображенного на рис. 2, к правому с помощью предыдущего канала.

него) двоичный симметричный канал с вероятностью p_2 , такой, что $p_1 = pp_2 + qq_2$, то полное устройство будет действовать как второй канал на рис. 2. Практически требуемое статистическое устройство можно осуществить с помощью соответствующей схемы со случайным элементом. Приведенные примеры наводят на мысль считать по определению, что канал K_1 с матрицей переходных вероятностей $\|p_i(j)\|$ включает канал K_2 с матрицей переходных вероятностей $\|q_i(j)\|$, если существуют стохастические матрицы A и B , такие, что

$$A \| p_i(j) \| = B \| q_i(j) \|.$$

Такое определение возможно, но оно допускает в действительности некоторое обобщение, при котором еще сохраняются те свойства, которые желательно получить и для отношения включения каналов. А именно, можно рассматривать «предварительную» и «заключительную» статистические операции, коррелированные между собой. Физически можно представить себе устройства, расположенные на передающем и приемном конце канала и содержащие случайные, но не обязательно независимые элементы. Например, статистические устройства могут производить выбор с лент, между которыми имеется некоторая корреляция. Ясно, что такая ситуация вполне реальна. Математически она может быть представлена в простейшем случае следующим образом.

Определение. Пусть $p_i(j)$, ($i = 1, \dots, a$; $j = 1, \dots, b$) есть вероятности перехода для дискретного канала без памяти K_1 и $q_k(l)$ ($k = 1, \dots, c$; $l = 1, \dots, d$) есть такие же вероятности для канала K_2 . Будем говорить, что канал K_1 включает канал K_2 , т. е. $K_1 \supseteq K_2$, тогда и только тогда, когда существуют два множества вероятностей перехода $r_{\alpha h}(i)$ и $t_{\alpha j}(l)$, такие, что

$$r_{\alpha h}(i) \geq 0, \quad \sum_i r_{\alpha h}(i) = 1,$$

и

$$t_{\alpha j}(l) \geq 0, \quad \sum_l t_{\alpha j}(l) = 1,$$

и, кроме того, существуют

$$g_\alpha \geq 0, \quad \sum_\alpha g_\alpha = 1,$$

при которых выполняется равенство

$$\sum_{\alpha, i, j} g_\alpha r_{\alpha h}(i) p_i(j) t_{\alpha j}(l) = q_h(l). \quad (1)$$

Грубо говоря, за этим определением скрывается некоторое множество предшествующих и последующих каналов R_α и T_α для канала K_1 , которые используются попарно, и вероятность для такой пары равна g_α . Когда операции указанного вида производятся над каналом K_1 , он действует аналогично каналу K_2 .

Определим *чистый канал* требованием, чтобы все переходы имели вероятности либо 0, либо 1; иными словами, в таком канале каждая буква на входе безошибочно переходит в определенную букву на выходе. Любую пару, составленную из предварительного и заключительного каналов R_α и T_α , фигурирующую в (1), можно представить себе как взвешенную сумму чистых каналов, предварительного и заключительного для канала K_1 и действующих на него. Рассмотрим всевозможные отображения входных букв канала R_α на его выходные буквы и сопоставим каждому отображению вероятность,

такую, чтобы в результате получился канал, эквивалентный R_α . Тогда отображение, при котором каждая буква k отображается в некоторую букву m_k , имеет вероятность $\prod_k r_{\alpha k}(m_k)$. Подобная же

операция может быть проделана над заключительным каналом T_α и над парой из предварительного и заключительного каналов, причем вероятности берутся равными произведению вероятностей.

Указанное сведение к «чистым» компонентам может быть выполнено при каждом α с вероятностями, соответствующими данному α . Таким образом, операции, которые должны быть выполнены в соответствии с формулой (1), могут быть выполнены при помощи только чистых каналов на основании указанного выше сведения. Другими словами, отношение включения между каналами может быть эквивалентным образом переформулировано так, что все $r_{\alpha k}(i)$ и $t_{\alpha j}(l)$ будут соответствовать чистым каналам.

Отношение включения транзитивно. Если $K_1 \supseteq K_2$ и $K_2 \supseteq K_3$, то $K_1 \supseteq K_3$. В самом деле, пусть g_α , R_α , T_α суть вероятности для предварительного и заключительного каналов для первого включения и g_β , R_β , T_β — аналогичные характеристики для второго. Тогда при помощи вероятностей $g_\alpha g_\beta$ для предварительного канала $R_\beta \cup R_\alpha$ и для заключительного канала $T_\alpha \cup T_\beta$ (знак $\cup$ означает последовательное использование каналов, соответствующее произведению матриц) получим канал K_3 из канала K_1 . Если $K_1 \supseteq K_2$ и $K_2 \supseteq K_1$, то будем говорить, что каналы K_1 и K_2 эквивалентны. Обозначим это через $K_1 \equiv K_2$. Заметим, что всегда $K_1 \equiv K_2$. Группируя каналы в классы эквивалентности, получаем частичное упорядочение дискретных каналов без памяти. Очевидно, канал с одной входной буквой, с одной выходной буквой и с вероятностью 1 для перехода является общей нижней границей для всех каналов. Общей (конечной) верхней границы не существует. Однако если ограничиться каналами, у которых не больше n входных и выходных символов (или им эквивалентными), то для этого подмножества можно найти верхнюю границу, а именно чистый канал с n входными и n выходными буквами, взаимно однозначно отображаемыми друг на друга.

Отношение включения сохраняется при сложении и умножении каналов. Если K_1, K'_1, K_2, K'_2 — каналы и $K_1 \supseteq K'_1$ и $K_2 \supseteq K'_2$, то

$$K_1 + K_2 \supseteq K'_1 + K'_2, \quad K_1 K_2 \supseteq K'_1 K'_2.$$

Напомним, что употребляемые здесь понятия суммы и произведения каналов были определены в более ранней работе автора¹⁾ как такие каналы, в которых может использоваться либо канал K_1 , либо

¹⁾ Шеннон К., Пропускная способность канала с шумом при нулевой ошибке (см. данный сборник, стр. 464.—Прим. ред.).

K_2 (сумма каналов), или как такие, в которых одновременно используются каналы K_1 и K_2 (произведение каналов). Для доказательства выписанных неравенств для произведения предположим, что канал K_1' получается из канала K_1 с помощью $(g_\alpha, R_\alpha, T_\alpha)$ и канал K_2' из канала K_2 с помощью $(g_\beta, R_\beta, T_\beta)$. Тогда с помощью $(g_\alpha g_\beta, R_\alpha R_\beta, T_\alpha T_\beta)$ канал $K_1' K_2'$ получается из канала $K_1 K_2$, где произведение $R_\alpha R_\beta$ означает произведение каналов. Для суммы каналов доказательство аналогично. Именно сумма $K_1' + K_2'$ получается из $K_1 + K_2$ с помощью $(g_\alpha g_\beta, R_\alpha + R_\beta, T_\alpha + T_\beta)$, где знак «+» означает сумму каналов, а индексы α, β пробегают по всевозможным значениям.

Если в дискретном канале без памяти рассмотреть слова из n букв, то получим другой канал без памяти, в котором входными буквами будут служить n -буквенные слова, составленные из выходных букв нашего канала. Очевидно, что такой канал эквивалентен K^n . Следовательно, если $K_1 \supseteq K_2$, то канал K_1^n включает канал K_2^n .

Предположим, что $K_1 \supseteq K_2$ и $K_1 \supseteq K_3$ и что K_1 и K_3 имеют матрицы $\|p_i(j)\|$ и $\|q_i(j)\|$ соответственно. Тогда K_1 включает канал, матрица которого имеет вид

$$\lambda \|p_i(j)\| + (1 - \lambda) \|q_i(j)\| \quad (0 \leq \lambda \leq 1).$$

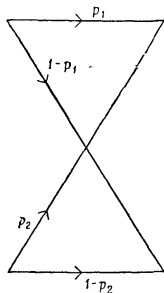
Следовательно, если каналы характеризовать матрицами перехода, то совокупность каналов, меньших данного канала K_1 , образует в пространстве матриц выпуклое тело. В самом деле, канал с матрицей $\lambda \|p_i(j)\| + (1 - \lambda) \|q_i(j)\|$ может быть получен из канала K_1 объединением $(\lambda g_\alpha, R_\alpha, T_\alpha)$ и $((1 - \lambda) g_\beta, R_\beta, T_\beta)$.

Наш основной результат, явившийся главной причиной для рассмотрения отношения включения каналов, связывает высказанные выше соображения с теорией кодирования. Покажем, что для всякого кода для канала K_2 и любого канала $K_1 \supseteq K_2$ найдется код для канала K_1 , не худший, чем код для канала K_2 в смысле величины вероятности ошибки.

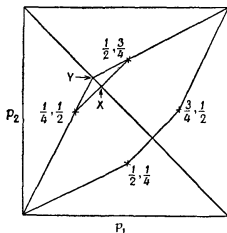
Теорема. Предположим, что $K_1 \supseteq K_2$ и что для канала K_2 задано множество кодовых слов длины n : $W_1, W_2, \dots, W_m$ и система декодирования такова, что если слово W_i употребляется с вероятностью P_i , то средняя вероятность ошибки равна P_e . Тогда для канала K_1 существует множество t кодовых слов длины n и такая декодирующая система, что если эти слова употреблять с теми же вероятностями P_i , то средняя вероятность ошибки $P'_e < P_e$. Следовательно, пропускная способность канала K_1 не меньше, чем пропускная способность канала K_2 .

Доказательство. Пусть $K_1 \supseteq K_2$ и множество $(g_\alpha, R_\alpha, T_\alpha)$ превращает канал K_1 в канал K_2 , где T_α и R_α — чистые каналы.

При каждом α канал R_α определяет отображение входных слов из кода K_2^n в множество входных слов из словаря канала K_1^n (т. е. слов, в которые R_α преобразует вводимый код). Далее T_α определяет отображение выходных букв канала K_1 на выходные буквы канала K_2 . Из системы кодирования и декодирования для канала K_2 можно получить для каждого значения α кодирующую и декодирующую систему для канала K_1 . Возьмем для этого в качестве кода слова,



Р и с. 4. Общий двонч-
ный канал.



Р и с. 5. Шестиугольник из двоич-
ных каналов, включенных в ти-
пичный двоичный канал.

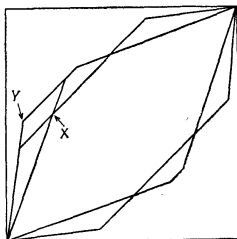
полученные при отображении, задаваемом R_α , из кодовых слов для K_2 . Декодирование определим следующим образом: слово, входящее из канала K_1 , декодируется так же, как декодируется слово, которое получается из данного при отображении T_α . Такой код дает среднюю вероятность ошибки $P_{e\alpha}$ для канала K_1 . Теперь очевидно, что

$$P_e = \sum_a g_a P_{e\alpha}, \quad (2)$$

поскольку канал K_2 действует как бы с построенными нами различными кодами с вероятностью g_α каждый. Так как из последнего уравнения вытекает, что взвешенное среднее вероятностей $P_{e\alpha}$ равно P_e , то должно существовать по крайней мере одно такое $P_{e\alpha}$, которое больше или равно P_e . [Если все $P_{e\alpha}$ были бы меньше P_e , то правая часть (2) необходимо была бы меньше P_e .] Система кодирования и декодирования для этого частного значения α и дает доказательство теоремы. В самом деле, величина $P_{e \text{ opt}}(M, n)$, равная мини-

муму вероятности ошибки при передаче M равновероятных слов длины n , будет для канала K_1 не больше, чем для канала K_2 . Аналогичным образом пропускная способность канала, определяемая как наибольшая нижняя граница скоростей передачи, для которых P_e можно заставить стремиться к нулю, будет для канала K_1 по крайней мере столь же большой, как и для канала K_2 .

Интересно проследить геометрическую интерпретацию отношений включения каналов в том простом случае, когда число входных и выходных букв равно двум (общий двоичный канал). Такой канал определяется двумя вероятностями p_1 и p_2 (рис. 4) и может быть представлен точкой единичного квадрата. В связи с этим см. работу Силвермана¹⁾, в которой пропускная способность канала и другие его параметры были изображены линиями на таком квадрате. На рис. 5 канал, для которого $p_1 = 1/4$, $p_2 = 1/4$, изображен вместе с тремя другими, эквивалентными ему каналами с вероятностями p_2, p_1 ; $1 - p_2, 1 - p_1$; $1 - p_1, 1 - p_2$. Добавляя



Р и с. 6. Наибольшая нижняя граница X и наименьшая верхняя граница Y двух сравниваемых двоичных каналов.

две точки $(0,0)$ и $(1,1)$, получим шесть точек. Внутри шестиугольника, определяемого ими, заключены все каналы, включаемые в данный. В самом деле, непосредственной проверкой легко убедиться, что для каждого канала внутри этого шестиугольника найдется предварительный и заключающий каналы, которые приведут к одному из шести выделенных каналов. Очевидно также, что любая взвешенная сумма каналов с вероятностями g_a не выведет нас за пределы шестиугольника.

¹⁾ Silverman R., On binary channels and their cascades, *IRE Trans. Inform. Theory*, IT-1 (1955), 190.

На рис. 5 двоичные симметричные каналы лежат на диагонали, идущей от $(1,0)$ к $(0,1)$. Из этого следует, что взятый нами канал включает, в частности, некоторый двоичный симметричный канал X , изображенный точкой $1/2 [p_1 + (1 - p_2)]$, $1/2 [p_2 + (1 - p_1)]$ квадрата, что для нашего примера дает $(3/5, 5/8)$. Кроме того, наш канал включается в двоичный симметричный канал Y с координатами $[p_1/(p_1 + p_2), p_2/(p_1 + p_2)]$, что для нашего частного случая превращается в $1/3, 2/3$. Такие включения дают простые верхние и нижние оценки для пропускной способности общего двоичного канала через более удобные характеристики подходящим образом подобранных симметричных каналов.

Случай, когда имеется два канала, ни один из которых не включает другого, представлен на рис. 6 с помощью двух шестиугольников. Здесь существует наибольшая нижняя граница и наименьшая верхняя граница для этих двух каналов, представленная на рис. 6 точками X и Y соответственно. Таким образом, в случае двоичных каналов имеется больше, чем простое частичное упорядочение, т. е. имеется структура.

Дальнейшее обобщение и некоторые предположения

Нам не удалось установить, определяет ли введенное здесь частичное упорядочение структуру в случае канала с n буквами. Множество точек, включающихся в данный канал, может быть найдено так же, как в случае, представленном на рис. 5, т. е. как выпуклая оболочка точек, полученных из канала при помощи предварительных и заключительных чистых каналов, но остается, например, неясным, отвечает ли пересечение двух таких выпуклых множеств какому-либо каналу.

Другой вопрос относится к обращению доказанной выше теоремы о кодах. Верно ли, что упорядочение, которое было определено выше, является в некотором смысле наиболее общим из таких, для которых справедлива доказанная теорема ¹⁾?

Понятие включения каналов может быть обобщено многими способами на каналы с памятью; следует отметить, что в другой работе ²⁾ использовалось такого рода понятие, только на более простом уровне для получения некоторых результатов в теории кодирования. Однако не ясно, каким должно быть наиболее естественное обобщение, пригодное в этих более общих случаях.

¹⁾ Shannon C., Zero error capacity of a noisy channel. (Русский перевод см. стр. 464 данного сборника.—Прим. ред.). (Эта проблема была решена в отрицательном смысле М. Л. Кармазиным. Его работа будет опубликована в сборнике «Проблемы математики».—Прим. ред.)

²⁾ Shannon C., Certain results in coding theory for noisy channel. (Русский перевод см. стр. 509 данного сборника.—Прим. ред.)

ВЕРОЯТНОСТЬ ОШИБКИ ДЛЯ ОПТИМАЛЬНЫХ КОДОВ В ГАУССОВСКОМ КАНАЛЕ¹⁾

Краткое содержание

В работе изучены кодирующие и декодирующие системы для непрерывного канала с аддитивным гауссовским шумом и ограничением средней мощности на передатчике. Определены верхняя и нижняя границы вероятности ошибки декодирования для оптимальных кодов и систем декодирования. Эти границы сближаются для скоростей передачи сигналов, близких к пропускной способности канала, а также для скорости передачи, близкой к нулю, но расходятся в промежуточных случаях. Приведены кривые, описывающие эти границы.

1. Введение

Рассмотрим канал связи следующего типа: один раз в секунду на передающем конце выбирается некоторое действительное число. Это число передается на приемный конец, искажаясь под действием аддитивного гауссовского шума, так что i -е действительное число s_i принимается как $s_i + x_i$. Здесь x_i предполагаются независимыми гауссовскими случайными величинами с одной и той же дисперсией N .

Кодовое слово длины n для такого канала — это последовательность из n действительных чисел $(s_1, s_2, \dots, s_n)$. Оно может быть геометрически представлено точкой n -мерного евклидова пространства. Эффект воздействия шума при этом будет состоять в перемещении этой точки в близлежащую в соответствии со сферическим гауссовским распределением.

Блочный код длины n с M словами есть преобразование целых чисел $1, 2, \dots, M$ во множество M кодовых слов $\omega_1, \omega_2, \dots, \omega_M$ (все они не обязательно разные). Геометрически блочный код представляется набором M (или менее) точек с соотношенными им целыми числами. Его можно рассматривать как способ передачи целых чисел от 1 до M на приемный конец (путем посылки соответствующего кодового слова). Декодирующая система для такого кода есть разбиение

¹⁾ Shannon C., Probability of error for optimal codes in a gaussian channel, *Bell System Techn. J.*, 38, May (1959), 611.

n -мерного пространства на M подмножеств, соответствующих целым числам от 1 до M . Это есть способ решения на приемном конце, какое целое число было передано. Если принятый сигнал попал в подмножество S_i , то принимается, что передаваемое сообщение соответствует целому числу i .

Предположим, что все целые числа от 1 до M соответствуют сообщениям, имеющим одну и ту же вероятность $1/M$. Для заданных кодирующей и декодирующей систем имеется определенная вероятность ошибки при передаче сообщения. Она дается формулой

$$P_e = \frac{1}{M} \sum_{i=1}^M P_{ei},$$

в которой P_{ei} есть вероятность того, что при передаче кодового слова ω_i оно декодируется некоторым целым числом, отличным от i . Очевидно, P_{ei} есть полная вероятность, определяемая гауссовским распределением с центром в ω_i в области, дополнительной к S_i .

Оптимальная декодирующая система для некоторого кода — это такая система, которая минимизирует вероятность ошибки для данного кода. Поскольку гауссовская плотность вероятности является монотонно убывающей функцией расстояния, оптимальная декодирующая система для некоторого данного кода это та, которая декодирует принятый сигнал целым числом, соответствующим кодовому слову, геометрически ближайшему к сигналу. Если имеется несколько кодовых слов с одним и тем же минимальным расстоянием до сигнала, то можно принять любое из них в качестве переданного; это не изменит вероятность ошибки. Декодирующая система такого вида называется системой декодирования по *минимуму расстояния* или *максимуму правдоподобия*. Она приводит к разбиению n -мерного пространства на n -мерные полиэдры или политопы вокруг различных сигнальных точек; каждый полиэдр ограничен конечным (не большим чем $M - 1$) числом $(n - 1)$ -мерных гиперплоскостей.

Нас интересует задача отыскания хорошего кода, т. е. такого расположения M точек, при котором минимизируется вероятность ошибки P_e . Если не накладывать ограничений на кодовые слова, то очевидно, что вероятность ошибки может быть сколь угодно малой при любых M , n и N , если разместить кодовые слова в n -мерном пространстве на достаточно большом расстоянии друг от друга. Однако в практических задачах на выбор кодовых слов накладываются обычно ограничения, не позволяющие использовать этот метод. Интересным случаем, рассматривавшимся прежде, является тот, когда некоторым образом ограничена средняя мощность кодовых слов, т. е. расстояния точек от начала координат не должны быть слишком большими.

Можно определить три различных возможных ограничения этого типа:

1. Все кодовые слова имеют *точно одну и ту же мощность* P , или одно и то же расстояние от начала координат. Таким образом, мы требуем, чтобы в качестве кодовых слов были выбраны точки, расположенные на поверхности сферы радиуса $\sqrt{nP}$.

2. Все кодовые слова имеют мощность P или меньшую. В этом случае все кодовые слова должны находиться внутри или на поверхности сферы радиуса $\sqrt{nP}$.

3. Средняя мощность всех кодовых слов равна P или меньше. В этом случае отдельные кодовые слова могут иметь квадрат расстояния от начала, больший чем nP , но совокупность квадратов расстояния не может превосходить nP .

Эти три случая, как будет показано, приводят к очень сходным результатам. Первое условие наиболее просто и приводит к несколько более отчетливым выводам. Сначала будет проанализирован именно этот случай, затем полученные результаты будут использованы для двух других. Поэтому *до тех пор, пока не оговаривается обратное, принимается, что все кодовые слова расположены на поверхности сферы радиуса $\sqrt{nP}$* .

Первая задача состоит в том, чтобы по возможности точно оценить вероятность ошибки $P_e(M, n, \sqrt{P/N})$ для наилучшего кода длины n , содержащего M слов, каждое мощностью P , которые искажаются шумом с дисперсией N . Эту минимальную или *оптимальную* вероятность ошибки обозначим через $P_{e \text{ опт}}(M, n, \sqrt{P/N})$. Ясно, что для фиксированных M и n , $P_{e \text{ опт}}$ будет являться функцией лишь величины $A = \sqrt{P/N}$ при изменении масштаба геометрического представления. Получим нижнюю и верхнюю границы $P_{e \text{ опт}}$ нескольких разных типов. В важной области значений эти границы достаточно близки друг к другу и дают хорошие оценки $P_{e \text{ опт}}$. Дадим некоторые вычисленные значения и кривые; найденные границы служат для определения аналогичных границ для второго и третьего типов условий, налагаемых на кодовые слова.

Геометрический подход используется так же, как это делалось автором и ранее¹⁾. Однако теперь все доведено до численных результатов. Рассмотренная задача близка к задаче, изучавшейся Райсом²⁾, получившим оценку, подобную одной из наших верхних оценок, но не столь точную. Границы, аналогичные найденным здесь,

¹⁾ Shannon C., Communication in the presence of noise. (Русский перевод см. стр. 433 данного сборника.—Прим. ред.)

²⁾ Rice S. O., Communication in the presence of noise—probability of error for two encoding schemes, *BSTJ*, 29, January (1950), 60.

приведены Элайесом¹⁾ для случая двоичного симметричного и двоичного стирающего каналов²⁾. Эти результаты связаны также с границами, найденными автором для общего случая дискретного канала без памяти³⁾.

Обе найденные здесь границы, нижняя и верхняя, имеют в общих чертах экспоненциальную зависимость от n при фиксированной скорости передачи R и фиксированном значении P/N . Они могут быть представлены (при условии, что $R = (1/n) \log M$, так что R является скоростью передачи для данного кода) в форме

$$e^{-E(R)n+o(n)}, \quad (1)$$

где $E(R)$ — некоторая функция R (а также P/N , рассматриваемого как фиксированный параметр). В выражении (1) $o(n)$ является членом меньшего порядка, чем n , так что при $n \rightarrow \infty$ он становится малым по сравнению с $E(R)n$.

Таким образом, при больших n логарифм величины границы линейно убывает с ростом n , или, точнее, отношение этого логарифма к n приближается к постоянному значению $E(R)$. Величина $E(R)$ является грубой мерой того, как быстро вероятность ошибки стремится к нулю. Назовем величину такого типа *надежностью*. Точнее, можно определить надежность для некоторого канала следующим образом:

$$E(R) = \lim_{n \rightarrow \infty} \left[-\frac{1}{n} \log P_{e \text{ opt}}(R, n) \right], \quad (2)$$

где $P_{e \text{ opt}}(R, n)$ — оптимальная вероятность ошибки для кода длины n при скорости передачи R . При этом обнаружим, что границы *точно* определяют $E(R)$ в некоторой важной области значений скорости передачи, начиная от некоторой критической скорости передачи R_c и до значения пропускной способности канала. Между нулем и R_c надежность E не определяется найденными границами точно, но находится в пределах не слишком широкого интервала.

Относительно надежности E нужно заметить, что в вышеописанном случае знание $E(R)$ и n в выражении (1) не дает возможности точно определить вероятность ошибки даже при больших n ; член $o(n)$ может обусловить появление большого и растущего множителя. С другой стороны, если задана вероятность ошибки и значение $E(R)$, то необходимую длину кода n можно *точно* опре-

¹⁾ Э л а й е с П., Кодирование для двух каналов с шумами, в сб. «Теория передачи сообщений», М., ИЛ, 1957, 114.

²⁾ См. также работу: Д о б р у ш и н Р. Л., Асимптотические оценки вероятности ошибки при передаче сообщения по дискретному каналу связи без памяти с симметрической матрицей вероятностей перехода, *Теория вероятностей и ее прил.*, VII, 3 (1962) 283—311. — *Прим. ред.*

³⁾ S h a n n о n С., Certain results in coding theory for noisy channels. (Русский перевод см. стр. 509 данного сборника. — *Прим. ред.*)

делить при большом n , так как n асимптотически равно величине $-(1/E) \log P_e$. Эта обратная задача, возможно, практически наиболее обычна: задан допустимый уровень вероятности ошибки, спрашивается, какова должна быть длина кода?

Рассматриваемый здесь тип канала тесно связан с каналом, имеющим ограниченную полосу пропускания (W гц), находящимся под воздействием белого гауссовского шума. В известном смысле такой ограниченный по полосе канал можно рассматривать как канал, имеющий $2W$ координат в 1 сек.; каждая из координат независимо искажается гауссовской случайной величиной. Однако такое отождествление нужно применять с осторожностью, так как для того, чтобы контролировать физически все эти степени свободы, оставаясь строго внутри заданной полосы частот, потребуется бесконечная задержка.

Можно очень точно уложиться в пределы полосы W при большом, но конечном времени T , применяя, например, импульсы $(\sin x)/x$, у которых усекается хвост на расстоянии T от точки максимума. Такое усечение вызывает выход за пределы полосы с мощностью, не большей, чем мощность усеченной части, величина которой менее $1/T$ для мощности, имеющей место в случае $\sin x$. Увеличивая T , можно приблизиться к случаю, когда при ограниченной полосе, а также нулевой вероятности ошибки скорость передачи будет близкой к пропускной способности канала.

Однако для рассматриваемых здесь задач отрезок времени, связанный с вероятностью ошибки, играет большую роль и в случае применения полученных результатов к каналам, ограниченным по полосе, нужно считаться с дополнительным временем, необходимым для того, чтобы оставаться в пределах заданной полосы частот. В этом и заключается причина, по которой введено принятое в данной работе определение канала.

2. Сводка результатов

Ниже вкратце изложены результаты, содержащиеся в работе как для облегчения ее чтения, так и для удобства читателей, которые могут интересоваться результатами и не желают разбираться в проводимом ниже детальном анализе. Нужно заметить, что алгебраические выкладки, содержащиеся в некоторых местах работы, необычайно громоздки. Введем следующие обозначения:

P — мощность сигнала (все кодовые слова расположены на поверхности сферы радиуса $\sqrt{nP}$);

N — мощность шума (N является дисперсией каждой составляющей шума);

$A = \sqrt{P/N}$ — отношение сигнал/шум по амплитуде;

n — число измерений или блоковая длина кода;
 M — число кодовых слов;
 $R = (1/n) \log M$ — скорость передачи кодом (в натуральных единицах);
 $C = 1/2 \log (P+N) / N = 1/2 \log (A^2+1)$
 — пропускная способность канала (на степень свободы);

θ — переменная, соответствующая половине угла раствора (полууглу) конуса, встречающаяся в геометрических задачах, приведенных ниже;

$\Omega(\theta)$ — телесный угол (в n -мерном пространстве) конуса с полууглом θ или площадь, вырезаемая на единичной n -мерной сфере конусом;

$\theta_0 = \arcsin A$ — угол конуса, соответствующий пропускной способности канала;

θ_1 — угол конуса такой, что телесный угол сферы равен $\Omega(\pi)$; таким образом, θ_1 — угол конуса, соответствующий скорости передачи R ;

$$G = G(\theta) = 1/2(A \cos \theta + \sqrt{A^2 \cos^2 \theta + 4})$$

— величина, часто встречающаяся далее в формулах;

θ_c — решение уравнения $2 \cos \theta_c - A G(\theta_c) \sin^2 \theta_c = 0$ (это критический угол, играющий большую роль, так как границы изменяются в зависимости от того,

$$\theta_1 > \theta_c \text{ или же } \theta_1 < \theta_c;$$

$Q(\theta) = Q(\theta, A, n)$ — вероятность того, что точка X n -мерного пространства с расстояния $A\sqrt{n}$ от начала координат попадет в область вне конуса с полууглом θ между вертикалью из начала координат O и осью OX (смещение определяется сферическим гауссовым распределением с единичной дисперсией по каждой координате);

$$E_L(\theta) = A^2/2 - 1/2 A G \cos \theta - \log (G \sin \theta)$$

— показатель в экспоненте, встречающийся в выражениях, определяющих значение границ;

$P_{e \text{ opt}}(n, R, A)$ — вероятность ошибки для лучшего кода длины n при отношении сигнал/шум A и скорости передачи R ;

$\varphi(X)$ — нормальное распределение с нулевым средним и единичной дисперсией.

Суммируем теперь результаты работы. Границы, в которых заключено значение $P_{e \text{ opt}}$, можно выразить следующим образом:

$$Q(\theta_1) \leq P_{e \text{ opt}} \leq Q(\theta) - \int_0^{\theta_1} \frac{\Omega(\theta)}{\Omega(\theta_1)} dQ(\theta) \quad (3)$$

(здесь $dQ(\theta)$ отрицательно, так что последний член в правой части неравенства положителен). Эти границы можно записать в форме довольно сложных интегралов. Для того чтобы достаточно полно изучить их поведение, дадим, во-первых, асимптотические выражения для них при больших n и, во-вторых, более грубые значения, которые, однако, выражаются с помощью элементарных функций без интегралов.

Асимптотическое выражение нижней границы (асимптотика справедлива при $n \rightarrow \infty$) имеет вид

$$Q(\theta_1) \sim \frac{1}{\sqrt{n\pi} G \sqrt{1+G^2} \sin \theta_1 (\cos \theta_1 - AG \sin^2 \theta_1)} e^{-E_L(\theta_1)n} = \\ = \frac{\alpha(\theta_1)}{\sqrt{n}} e^{-E_L(\theta_1)n} \quad (\theta_1 > \theta_0), \quad (4)$$

асимптотическое выражение верхней границы —

$$Q(\theta_1) - \int_0^{\theta_1} \frac{\Omega(\theta)}{\Omega(\theta_1)} dQ(\theta) \sim \frac{\alpha(\theta_1)}{\sqrt{n}} e^{-E_L(\theta_1)n} \left(1 - \frac{\cos \theta_1 - AG \sin^2 \theta_1}{2 \cos \theta_1 - AG \sin^2 \theta_1} \right). \quad (5)$$

Эти формулы справедливы для $\theta_0 < \theta < \theta_c$. В этой области асимптотические выражения для нижней и верхней границ отличаются лишь множителем в скобках, не зависящим от n . Таким образом, *вероятность ошибки определена этими соотношениями асимптотически с точностью до множителя, зависящего от скорости передачи*. Для скоростей передач, близких к пропускной способности канала (θ_1 близко к θ_0), этот множитель лишь немного больше единицы и границы близки друг к другу. Для более низких скоростей передачи, близких к R_c (соответствующей θ_c), множитель становится большим. Для $\theta_1 > \theta_c$ верхняя граница асимптотически равна

$$\frac{1}{\cos \theta_c \sin^3 \theta_c G(\theta_c) \sqrt{\pi E''(\theta_c) [1+G(\theta_c)]^2}} e^{-n[E_L(\theta_c)-R]}. \quad (6)$$

Кроме асимптотической границы, также дадим точные границы, справедливые для всех значений n ; однако они хуже асимптотических границ при больших n . Точная нижняя граница выражается соотношением

$$P_e \geq \frac{1}{6} \frac{\sqrt{n-1} e^{3/2}}{n(A+1)^3 e^{(A+1)^2/2}} e^{-E_L(\theta_1)n}. \quad (7)$$

Легко заметить, что точная граница равна асимптотической границе, умноженной на множитель, по существу не зависящий от n . Точная верхняя граница {справедливая, если максимум $G^n (\sin \theta)^{2n-3} \exp[-(n/2)(A^2 - AG \cos \theta)]$ в области между 0 и θ_1 имеет место при θ_1 } выражается формулой

$$P_{e \text{ opt}} \leq \theta_1 \sqrt{2n} e^{3/2} G^n(\theta_1) \sin \theta_1^{n-2} \exp \left[\frac{n}{2} (-A^2 + AG \cos \theta_1) \right] \times \\ \times \left\{ 1 + \frac{1}{n \theta_1 \min [A, AG(\theta_1) \sin \theta_1 - \operatorname{ctg} \theta_1]} \right\}. \quad (8)$$

Для скоростей передачи, близких к пропускной способности канала, нижняя и верхняя асимптотические границы сходятся друг к другу, так что при большом n и малом $C-R$ (положительном)

$$P_{e \text{ opt}} \approx \Phi \left[\sqrt{n} \sqrt{\frac{2P(P+N)}{N(P+2N)}} (R-C) \right], \quad (9)$$

где Φ —нормальная функция распределения с единичной дисперсией.

Связь угла θ_1 в предыдущих формулах со скоростью передачи R определяется неравенствами

$$\frac{\Gamma \left(\frac{n}{2} + 1 \right) (\sin \theta_1)^{n-1}}{n \Gamma \left(\frac{n+1}{2} \right) \pi^{1/2} \cos \theta_1} \left(1 - \frac{1}{n} \operatorname{tg}^2 \theta_1 \right) \leq e^{-nR} \leq \\ \leq \frac{\Gamma \left(\frac{n}{2} + 1 \right) (\sin \theta_1)^{n-1}}{n \Gamma \left(\frac{n+1}{2} \right) \pi^{1/2} \cos \theta_1}. \quad (10)$$

Асимптотически они приводят к соотношению

$$e^{-nR} \sim \frac{\sin^n \theta_1}{\sqrt{2\pi n} \sin \theta_1 \cos \theta_1}. \quad (11)$$

Для малых скоростей передачи (в частности, при $R < R_c$) найденные выше границы расходятся, давая мало информации. Два других метода рассуждений приводят к другим границам, полезным при малой скорости передачи. Для верхней границы получаем

$$P_{e \text{ opt}} \leq \frac{1}{\lambda A \sqrt{\pi n}} e^{n[R - (\lambda^2 A^2)/4]}, \quad (12)$$

где λ удовлетворяет соотношению $R = \left(1 - \frac{1}{n} \right) \log (\sin 2 \arcsin \lambda / \sqrt{2})$. Заметим, что когда $R \rightarrow 0$ и $\lambda \rightarrow 1$, верхняя граница равна примерно

$1/(A\sqrt{\pi n}e^{-nA^{2/2}})$. Для нижней границы при малой скорости передачи получается

$$P_{e\text{opt}} \geq \frac{1}{2} \Phi \left[-A \left(\frac{2M}{(2M-1)} \frac{n}{2} \right)^{1/2} \right]. \quad (13)$$

Для больших M эта граница сходится к $\frac{1}{2} \Phi(-A\sqrt{n/2})$ и, если n велико, асимптотически равна $1/(A\sqrt{\pi n}e^{-nA^{2/2}})$. Таким образом, при скоростях передачи, близких к нулю, и при больших n , снова имеем случай, когда обе границы сходятся друг к другу и получается точное выражение для $P_{e\text{opt}}$.

Для кодов со скоростью передачи $R \geq C + \varepsilon$, где ε фиксировано и положительно, $P_{e\text{opt}}$ стремится к единице с возрастанием n .

3. Нижняя граница, основанная на соображениях «сферической упаковки»

Предположим, что имеется код, состоящий из M точек в n -мерном пространстве с расстояниями от начала координат, равными $\sqrt{nP}$. Так как любые два слова находятся на одном и том же расстоянии от начала координат, $(n-1)$ -мерная гиперплоскость, которая делит пополам соединяющий их отрезок и перпендикулярна к нему, проходит через начало координат. Таким образом, все гиперплоскости, определяющие полиэдры, окружающие эти точки (для оптимальной декодирующей системы), проходят через начало координат. Эти полиэдры поэтому имеют вид пирамид с вершинами в начале координат. Вероятность ошибки для такого кода составляет

$$\frac{1}{M} \sum_{i=1}^M P_{ei},$$

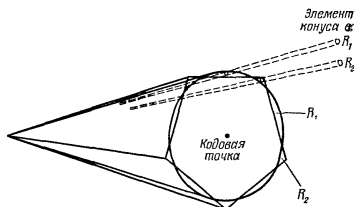
где P_{ei} равно вероятности того, что при использовании i -го кодового слова оно под воздействием шума окажется вне пирамиды, построенной вокруг i -го слова. Вероятность правильного приема равна

$$1 - \frac{1}{M} \sum_{i=1}^M P_{ei} = \frac{1}{M} \sum_{i=1}^M (1 - P_{ei}).$$

Это средняя вероятность того, что кодовое слово окажется в пределах своей пирамиды.

Пусть i -я пирамида имеет телесный угол Ω_i (т. е. Ω_i является площадью, вырезаемой пирамидой на единичной n -мерной сферической поверхности). Рассмотрим для сравнения правильный круговой n -мерный конус с тем же самым телесным углом Ω_i , имеющий кодовое слово на своей оси на расстоянии $\sqrt{nP}$ от начала координат.

Можно утверждать, что вероятность того, что кодовое слово останется в своем конусе, больше, чем вероятность того, что оно останется в своей пирамиде. Это следует из монотонного убывания плотности вероятностей с увеличением расстояния от кодового слова. Пирамиду можно деформировать в конус путем перемещения малых конических элементов с больших расстояний на меньшие расстояния от кодового слова. Такие перемещения непрерывно увеличивают вероятность. Для трехмерного случая это проиллюстрировано на рис. 1. Перемещение малых конических элементов из области вне конуса



Р и с. 1. Пирамида, деформируемая в конус перемещением малых элементов конуса с больших расстояний на меньшие.

в область внутри него увеличивает вероятность, так как плотность вероятности внутри конуса больше, чем вне его. Формально это можно показать интегрированием плотности вероятности по области R_1 конуса, лежащей вне пирамиды, и области R_2 пирамиды, лежащей вне конуса. Первая величина будет больше телесного угла Ω области R_1 , умноженного на плотность вероятности на границе конуса. Для пирамиды получится значение, меньшее этой величины.

Поэтому некоторая граница вероятности ошибки P_e для данного кода определится соотношением

$$P_e \geq \frac{1}{M} \sum_{i=1}^M Q^*(\Omega_i), \quad (14)$$

в котором Ω_i — телесный угол i -й пирамиды, а $Q^*(\Omega)$ — вероятность попадания точки во внешнюю область, окружающую конус с телесным углом Ω . Для телесного угла n -мерной сферы справедливо соотношение $\sum_{i=1}^M \Omega_i = \Omega_0$, так как наши пирамиды соответствуют разбиению сферы. Теперь, используя свойство плотности вероятности убывать с расстоянием, находим, что $Q^*(\Omega)$ является выпуклой

функцией Ω . Это позволяет произвести дальнейшее упрощение выражения для границы, заменив все Ω_i их средними Ω_0/M . Тогда

$$\frac{1}{M} \sum_{i=1}^M Q^*(\Omega_i) \geq Q^*\left(\frac{\Omega_0}{M}\right),$$

и поэтому

$$P_e \geq Q^*\left(\frac{\Omega_0}{M}\right).$$

Удобнее выразить границу вероятности ошибки через полуугол конуса θ , чем через телесный угол Ω . Определим $Q(\theta)$ как вероятность попадания точки в область вне конуса с полууглом θ . Тогда, если θ_1 соответствует конусу с телесным углом Ω_0/M , вышеприведенную границу можно представить в виде

$$P_e \geq Q(\theta_1). \quad (15)$$

Это и есть найденная *основная нижняя граница для P_e* . Необходимо выразить ее через P , N , M и дать оценку при помощи простых функций.

Следует заметить, что эта граница даст точную вероятность ошибки, которая имела бы место, если бы существовала возможность разбить все пространство на M конгруэнтных конусов, соответствующих передаваемым словам, и расположить кодовые слова на осях этих конусов. Интуитивно представляется весьма правдоподобным, что любой конкретный код должен иметь большую вероятность ошибки, чем код, основанный на таком коническом разбиении. Очевидно, что если $M > 2$, то такое разбиение можно легко сделать лишь для $n = 1$ или 2.

Нижнюю границу $Q(\theta_1)$ можно выразить в терминах известного в статистике нецентрального t -распределения¹⁾. Нецентральное t -распределение задает вероятность того, что отношение случайной величины $(z + \delta)$ к квадратному корню из среднего квадрата f других случайных величин

$$\sqrt{\frac{1}{f} \sum x_i^2}$$

не превосходит t , причем все переменные x_i и z являются независимыми гауссовскими случайными величинами с нулевыми средними и единичными дисперсиями, а δ постоянно. Обозначая эту

¹⁾ Johnson N. L. and Welch B. L., Applications of the noncentral t -distribution, *Biometrika*, 31 (1939), 362.

вероятность $P(f, \delta, t)$, имеем

$$P(f, \delta, t) = \text{Pr} \left\{ \frac{z + \delta}{\sqrt{\frac{1}{f} \sum_1^f x_i^2}} \leq t \right\}. \quad (16)$$

С геометрической точки зрения t -распределение соответствует сферическому гауссовскому распределению с единичной дисперсией вокруг точки с расстоянием δ от начала координат в $(f+1)$ -мерном пространстве. Вероятность $P(f, \delta, t)$ является вероятностью попадания во внешнюю область конуса, выходящего из начала координат и имеющего центр распределения на своей оси. Котангенс его полуугла θ равен $t/\sqrt{f}$, поэтому вероятность $Q(\theta)$ может быть задана формулой

$$Q(\theta) = P\left(n-1, \sqrt{\frac{nP}{N}}, \sqrt{n-1} \text{ctg } \theta\right). \quad (17)$$

Кажется, что нецентральное t -распределение не относится к распределениям, которые широко табулировались. Джонсон и Велч¹⁾ дают несколько таблиц, однако они предназначаются для приложений другого рода и в нашем случае пользование ими неудобно. Кроме того, они не доведены до больших значений n . Поэтому дадим оценку для нижней границы, выведя асимптотическую формулу для функции распределения $Q(\theta)$, а также ее плотности вероятности $dQ/d\theta$. Сначала, однако, найдем *верхнюю* границу $P_{\text{е опт}}$ в терминах того же распределения $Q(\theta)$.

4. Верхняя граница, определяемая методом случайного кодирования

Найдем верхнюю границу $P_{\text{е опт}}$, пользуясь соображениями, основанными на методе случайного кодирования. Рассмотрим ансамбль кодов, получающихся случайным размещением M точек на поверхности сферы радиуса $\sqrt{nP}$. Более точно, каждая точка размещается независимо от всех других с вероятностной мерой, пропорциональной площади поверхности или, что эквивалентно, телесному углу. Декодирование в каждом коде ансамбля осуществляется по правилу минимума расстояния.

Требуется вычислить *среднюю вероятность ошибки* для нашего ансамбля кодов.

Из-за симметрии расположения кодовых точек вероятность ошибки, усредненная по всему ансамблю, равна M -кратной средней вероят-

¹⁾ Johnson N. L., Welch B. L., Application of the noncentral t -distribution, *Biometrika*, 31 (1939), 362.

ности ошибки; соответствующей некоторой частной кодовой точке, например, кодовой точке 1. Вычислим ее следующим образом: вероятность того, что передано сообщение 1, равна $1/M$. Дифференциальная вероятность того, что из-за шума оно будет перемещено в область между конусами с полууглом θ и полууглом $\theta + d\theta$ (эти конусы имеют вершину в начале координат и общую ось — линию, соединяющую начало координат с точкой кодового слова 1), есть $-dQ(\theta)$. (Напомним, что по определению $Q(\theta)$ является вероятностью того, что из-за воздействия шума точка окажется вне конуса с углом θ и осью, проходящей через соответствующую точку передаваемого кодового слова и начало координат.) Рассмотрим теперь конус с полууглом θ , окружающий *точку принятого кодового слова* (но не конус, описанный вокруг соответствующей точки передаваемого слова, как это было выше). Если этот конус свободен от точек, соответствующих передаваемым кодовым словам, то принятое слово декодируется правильно как сообщение 1. Если он не свободен, то другие точки окажутся ближе и принятое слово будет декодировано неправильно. (Очевидно, что вероятность того, что две или более точки находятся точно на одном и том же расстоянии от принятого кодового слова, равна нулю и может не учитываться.)

Вероятность по ансамблю кодов того, что конус с полууглом будет пуст, можно легко вычислить. В самом деле, вероятность того, что конкретное кодовое слово, например кодовое слово 2 или 3 и т. д., попадет в конус, равна $\Omega(\theta)/\Omega(\pi)$, т. е. отношению телесного угла конуса ко всему телесному углу. Вероятность того, что конкретное кодовое слово не попадет в конус, равна $1 - \Omega(\theta)/\Omega(\pi)$. Вероятность того, что все другие $M - 1$ слов не попадут в конус, равна $[1 - \Omega(\theta)/\Omega(\pi)]^{M-1}$, так как в рассматриваемом ансамбле отдельные точки размещаются на поверхности сферы независимо друг от друга. Поэтому приращение вероятности ошибки за счет смещения точки 1 на угол в интервале $(\theta, \theta + d\theta)$ составляет величину $-1/M \{1 - [1 - \Omega(\theta)/\Omega(\pi)]^{M-1}\} dQ(\theta)$. Общая средняя вероятность по всем кодовым словам и всем шумовым смещениям равна

$$P_{er} = - \int_0^\pi \left\{ 1 - \left[1 - \frac{\Omega(\theta)}{\Omega(\pi)} \right]^{M-1} \right\} dQ(\theta). \quad (18)$$

Это есть точная формула для средней вероятности ошибки P_{er} нашего случайного ансамбля кодов. Поскольку это есть среднее значение от P_e для частных кодов, то должны существовать конкретные коды в этом ансамбле по крайней мере с такой же вероятностью ошибки, и поэтому $P_{e, \text{opt}} \leq P_{er}$.

Можно слегка ослабить полученные оценки, но зато получить более простые формулы следующим образом. Замечая, во-первых, что $\{1 - [\Omega(\theta)/\Omega(\pi)]^{M-1}\} \leq 1$, а также, используя хорошо известное

неравенство $(1-x)^n \geq 1-nx$, получаем $\{1 - [1 - \Omega(\theta)/\Omega(\pi)]^{M-1}\} \leq (M-1) [\Omega(\theta)/\Omega(\pi)] \leq M [\Omega(\theta)/\Omega(\pi)]$. Теперь разобьем интеграл на две части, интегрируя в пределах $0 \leq \theta \leq \theta_1$ и $\theta_1 \leq \theta \leq \pi$. В первой области используем только что приведенные неравенства, во второй области выражение в фигурных скобках заменим единицей. Тогда

$$P_{er} \leq - \int_0^{\theta_1} M \left[\frac{\Omega(\theta)}{\Omega(\pi)} \right] dQ(\theta) - \int_{\theta_1}^{\pi} dQ(\theta) \quad (19)$$

или

$$P_{er} \leq - \frac{M}{\Omega(\pi)} \int_0^{\theta_1} \Omega(\theta) dQ(\theta) + Q(\theta_1).$$

Для θ_1 удобно выбрать то же значение, которое появлялось при подсчете нижней границы, т. е. значение, при котором $\Omega(\theta_1)/\Omega(\pi) = 1/M$, другими словами, значение, при котором внутри конуса с полууглом θ_1 получается одна ожидаемая точка. Второй член в выражении (19) совпадает с нижней оценкой $P_{e \text{ opt}}$, найденной ранее. Итак, подводя итоги, имеем

$$Q(\theta_1) \leq P_{e \text{ opt}} \leq Q(\theta_1) - \frac{M}{\Omega(\pi)} \int_0^{\theta_1} \Omega(\theta) dQ(\theta), \quad (20)$$

где $M\Omega(\theta_1) = \Omega(\pi)$. Это и есть наши основные нижняя и верхняя границы для значений $P_{e \text{ opt}}$.

Теперь желательно вычислить и оценить значения $\Omega(\theta)$ и $Q(\theta)$.

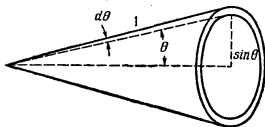
5. Формулы для скорости передачи R как функции угла конуса θ

Наши границы вероятности ошибки выражаются через угол конуса θ_1 , такой, что телесный угол этого конуса равен полному телесному углу сферы, умноженному на $1/M = e^{-nR}$. Чтобы связать эти величины более явно, найдем телесный угол n -мерного конуса с полууглом θ . Это означает вычисление площади $(n-1)$ -мерной поверхности, вырезаемой конусом на единичной сфере, как показано на рис. 2. Она получается суммированием вкладов от отдельных кольцевых элементов площади, заключенных между сферической $(n-1)$ -мерными поверхностями радиуса $\sin \theta$ и радиуса $\sin(\theta + d\theta)$. Таким образом, общая площадь «шапки» равна

$$\Omega(\theta) = \frac{(n-1) \pi^{(n-1)/2}}{\Gamma\left(\frac{n+1}{2}\right)} \int_0^{\theta_1} (\sin \theta)^{n-2} d\theta. \quad (21)$$

Здесь было использовано то обстоятельство, что площадь n -мерной сферы радиуса r задается формулой $S_n(r) = \pi^{n/2} r^{n-1} / \Gamma(n/2 + 1)$.

Чтобы получить простые неравенства и асимптотические выражения для $\Omega(\theta_1)$, произведем замену переменных в интеграле, введя $x = \sin \theta$, $d\theta = (1 - x^2)^{-1/2} dx$. Пусть $x_1 = \sin \theta_1$. Предположим,



Р и с. 2. «Шапка», вырезаемая из конуса единичной сферы.

что $\theta_1 < \pi/2$, так что $x_1 < 1$. Используя теорему о среднем значении, получим

$$(1 - x^2)^{-1/2} = (1 - x_1^2)^{-1/2} + \frac{\alpha}{(1 - \alpha^2)^{3/2}} (x - x_1), \quad (22)$$

где $0 \leq \alpha \leq x_1$. Слагаемое $\alpha(1 - \alpha^2)^{-3/2}$ должно лежать в пределах между 0 и $x_1(1 - x_1^2)^{-3/2}$, так как оно является монотонно возрастающей функцией. Поэтому имеем неравенства

$$(1 - x_1^2)^{-1/2} + \frac{(x - x_1) x_1}{(1 - x_1^2)^{3/2}} \leq (1 - x^2)^{-1/2} \leq (1 - x_1^2)^{-1/2} \quad (23)$$

при $0 \leq x \leq x_1$. Заметим, что $x - x_1$ отрицательно, так что поправочный член в левой части имеет нужный знак. Если воспользоваться этим в интеграле, выражающем $\Omega(\theta_1)$, получим

$$\begin{aligned} \frac{(n-1) \pi^{(n-1)/2}}{\Gamma\left(\frac{n+1}{2}\right)} \int_0^{x_1} x^{n-2} \left[(1 - x_1^2)^{-1/2} + \frac{(x - x_1) x_1}{(1 - x_1^2)^{3/2}} \right] dx &\leq \\ &\leq \Omega(\theta_1) \leq \frac{(n-1) \pi^{(n-1)/2}}{\Gamma\left(\frac{n+1}{2}\right)} \int_0^{x_1} x^{n-2} \frac{dx}{\sqrt{1 - x_1^2}}, \end{aligned} \quad (24)$$

$$\begin{aligned} \frac{(n-1) \pi^{(n-1)/2}}{\Gamma\left(\frac{n+1}{2}\right) \sqrt{1 - x_1^2}} \left[\frac{x_1^{n-1}}{n-1} + \frac{x_1^{n+1}}{n(1 - x_1^2)} - \frac{x_1^{n+1}}{(n-1)(1 - x_1^2)} \right] &\leq \\ &\leq \Omega(\theta_1) \leq \frac{(n-1) \pi^{(n-1)/2} x_1^{n-1}}{\Gamma\left(\frac{n+1}{2}\right) (n-1) \sqrt{1 - x_1^2}}, \end{aligned} \quad (25)$$

$$\frac{\pi^{(n-1)/2} (\sin \theta_1)^{n-1}}{\Gamma\left(\frac{n+1}{2}\right) \cos \theta_1} \left(1 - \frac{1}{n} \operatorname{tg}^2 \theta_1 \right) \leq \Omega(\theta_1) \leq \frac{\pi^{(n-1)/2} (\sin \theta_1)^{n-1}}{\Gamma\left(\frac{n+1}{2}\right) \cos \theta_1}. \quad (26)$$

Поэтому $\Omega(\theta_1)$ асимптотически совпадает с выражением в правой части при $n \rightarrow \infty$.

Площадь поверхности n -мерной сферы единичного радиуса равна $n\pi^{n/2}/\Gamma(n/2 + 1)$, поэтому

$$\frac{\Gamma\left(\frac{n}{2}+1\right)(\sin \theta_1)^{n-1}}{n\Gamma\left(\frac{n+1}{2}\right)\pi^{1/2}\cos \theta_1}\left(1-\frac{1}{n}\operatorname{tg}^2 \theta_1\right) \leq e^{-nR} =$$

$$= \frac{\Omega(\theta_1)}{\Omega(\pi)} \leq \frac{\Gamma\left(\frac{n}{2}+1\right)(\sin \theta_1)^{n-1}}{n\Gamma\left(\frac{n+1}{2}\right)\pi^{1/2}\cos \theta_1}. \quad (27)$$

Подставляя асимптотическое выражение для гамма-функции, получим

$$e^{-nR} = \frac{\sin^n \theta_1}{\sqrt{2\pi n} \sin \theta_1 \cos \theta_1} \left[1 + O\left(\frac{1}{n}\right)\right]. \quad (28)$$

Итак,

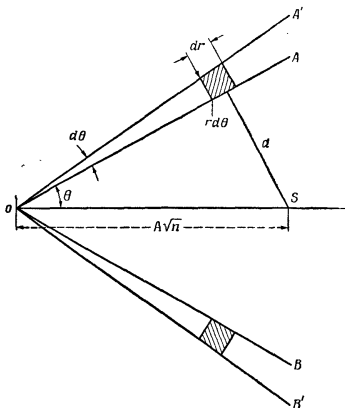
$$e^{-nR} \sim \sin^n \theta_1 / \sqrt{2\pi n} \sin \theta_1 \cos \theta_1 \text{ и } e^{-R} \sim \sin \theta_1.$$

При отыскании асимптотического значения для P_e необходимо применять более точное выражение для e^{-nR} , так как P_e изменяется в некоторое количество раз, когда θ_1 изменяется, например, на k/n . Однако, если нас интересует лишь надежность E , можно использовать простейшее соотношение $R \sim -\lg \sin \theta_1$.

6. Асимптотические формулы для $Q(\theta)$ и $Q'(\theta)$

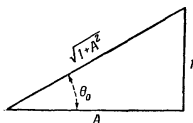
На рис. 3 O — начало координат, S — соответствующая точка кодового слова, а плоскость рисунка является плоским сечением n -мерного пространства. Линии OA и OB представляют собой конус (круговой) с углом θ вокруг OS (т. е. являются пересечением конуса с плоскостью рисунка). Линии OA' и OB' соответствуют несколько большему конусу с углом $\theta + d\theta$. Требуется оценить вероятность $-dQ_n(\theta)$ того, что точка S под действием шума попадет в область, лежащую между этими конусами. Далее, используя эту оценку, мы вычислим вероятность $Q_n(\theta)$ попадания точки S в область вне конуса с углом θ . Желательно в обоих случаях иметь асимптотическую оценку, простые формулы, отношение которых к истинному значению стремится к 1 при возрастании числа измерений n . Шум сдвигает независимо и нормально с дисперсией, равной 1, все координаты. Он приводит к сферическому гауссовому распределению в n -мерном пространстве. Плотность вероятности смещения точки на расстояние d равна

$$\frac{1}{(2\pi)^{n/2}} e^{-d^2/2} dV, \quad (29)$$



Р и с. 3. Сечение плоскостью конуса с полууглом раствора θ .

где dV — элемент объема. Прежде всего вычислим плотность вероятности (см. рис. 4) для заштрихованной области кругового кольца между двумя конусами и двумя сферами радиуса r и $r + dr$



Р и с. 4. Специальное значение θ_0 .

с центрами в начале координат. Расстояние до этого кольца из точки равно, согласно «теореме косинусов»,

$$d = (r^2 + A^2 n - 2rA\sqrt{n} \cos \theta)^{1/2}. \quad (30)$$

Дифференциальный объем области кругового кольца равен произведению $r dr d\theta$ и величины поверхности $(r-1)$ -мерной сферы радиуса $r \sin \theta$, т. е.

$$r dr d\theta \frac{(n-1)\pi^{(n-1)/2} (r \sin \theta)^{n-2}}{\Gamma\left(\frac{n+1}{2}\right)}. \quad (31)$$

Следовательно, дифференциальная вероятность попадания в область кругового кольца равна

$$\frac{1}{(\sqrt{2\pi})^n} \exp \left[\frac{-(r^2 + A^2 n - 2rA \sqrt{n} \cos \theta)}{2} \right] \times \\ \times \left[\frac{(n-1)\pi^{(n-1)/2} (r \sin \theta)^{n-2}}{\Gamma\left(\frac{n+1}{2}\right)} \right] r dr d\theta. \quad (32)$$

Дифференциальная вероятность $-dQ$ попадания в область между двумя конусами равна интегралу по r от этого выражения от нуля до бесконечности

$$-dQ = \frac{(n-1) d\theta}{2^{n/2} \sqrt{\pi} \Gamma\left(\frac{n+1}{2}\right)} \times \\ \times \int_0^\infty \exp \left[\frac{-(r^2 + A^2 n - 2rA \sqrt{n} \cos \theta)}{2} \right] (r \sin \theta)^{n-2} r dr. \quad (33)$$

В экспоненте можно представить $A^2 n$ в виде $A^2 n (\sin^2 \theta + \cos^2 \theta)$, тогда числитель дает полный квадрат

$$(r - A \sqrt{n} \cos \theta)^2,$$

а член $\sin^2$ можно вынести за знак интеграла. Получаем

$$-dQ = \frac{(n-1) \exp \left[-\frac{A^2 n \sin^2 \theta}{2} \right] (\sin \theta)^{n-2} d\theta}{2^{n/2} \sqrt{\pi} \Gamma\left(\frac{n+1}{2}\right)} \times \\ \times \int_0^\infty \exp \left[\frac{-(r - A \sqrt{n} \cos \theta)^2}{2} \right] r^{n-1} dr. \quad (34)$$

Теперь можно прямо приступить к оценке интеграла, который мы обозначим через K . Его можно точно выразить в виде конечной, но сложной суммы, включающей функции нормального распределения, возникающие при последовательном интегрировании по частям. Однако нас интересует получение простых асимптотических форм

для интеграла, когда $n \rightarrow \infty$. Эта задача полностью решена Дэвидом и Крускалом¹⁾, которые доказали в качестве леммы следующую асимптотическую формулу:

$$\int_0^{\infty} z^{\nu} \exp\left(-\frac{1}{2} z^2 + z \sqrt{\nu+1} \omega\right) dz \sim \sqrt{2\pi} \left(\frac{\bar{z}}{e}\right)^{\nu} \exp\left(\frac{1}{2} \bar{z}^2\right) T, \quad (35)$$

при следующих условиях: $\nu \rightarrow \infty$, ω фиксировано,

$$T = \left[1 + \frac{1}{4} (\sqrt{\omega^2 + 4} - \omega)^2\right]^{-1/2}$$

$$\text{и} \quad z = \frac{1}{2} \sqrt{\nu+1} \omega + \sqrt{\frac{1}{4} (\nu+1) \omega^2 + \nu}.$$

Доказательство основывается на том, что основной вклад в значение интеграла вносится в окрестности точки z , где подинтегральное выражение обращается в максимум. Вблизи этой точки, когда ν велико, подинтегральная функция ведет себя как нормальное распределение.

Интеграл K в формуле (34), который требуется вычислить, с точностью до множителя совпадает с интегралом леммы при $z = r$, $\omega = A \cos \theta$, $\nu = n - 1$. Интеграл приобретает вид

$$\begin{aligned} K &= \exp\left(-\frac{A^2 n \cos^2 \theta}{2}\right) \int_0^{\infty} z^{n-1} \exp\left[-\left(\frac{z^2}{2} + z A \sqrt{n} \cos \theta\right)\right] dz \sim \\ &\sim \exp\left(-\frac{A^2 n \cos^2 \theta}{2}\right) \sqrt{2\pi} \left(\frac{\bar{z}}{e}\right)^{n-1} T \exp\left(\frac{z^2}{2}\right). \end{aligned} \quad (36)$$

Имеем

$$\begin{aligned} z &= \frac{1}{2} \sqrt{n} A \cos \theta + \sqrt{\frac{1}{4} n A^2 \cos^2 \theta + n - 1} = \\ &= \sqrt{n} \left[\frac{1}{2} A \cos \theta + \sqrt{\frac{A^2}{4} \cos^2 \theta + 1 - \frac{1}{n}} \right] = \\ &= \sqrt{n} \left[\frac{1}{2} A \cos \theta + \sqrt{\frac{A^2}{4} \cos^2 \theta + 1} - \right. \\ &\quad \left. - \frac{1}{2n \sqrt{\frac{A^2}{4} \cos^2 \theta + 1}} + O\left(\frac{1}{n^2}\right) \right]. \end{aligned} \quad (37)$$

Подставляя

$$G = \frac{1}{2} [A \cos \theta + \sqrt{A^2 \cos^2 \theta + 4}],$$

¹⁾ David H., Kruskal W., The WAGR sequential t -test reaches a decision with probability one, *Ann. Math. Stat.*, 27, September (1956), 797.

имеем

$$\bar{z} = \sqrt{nG} \left[1 - \frac{1}{nG \sqrt{A^2 \cos^2 \theta + 4}} + O\left(\frac{1}{n^2}\right) \right],$$

так что

$$\begin{aligned} \left(\frac{\bar{z}}{e}\right)^{n-1} &= \left(\frac{\sqrt{nG}}{e}\right)^{n-1} \left[1 - \frac{1}{nG \sqrt{A^2 \cos^2 \theta + 4}} + O\left(\frac{1}{n^2}\right) \right]^{n-1} \sim \\ &\sim \left(\frac{\sqrt{nG}}{e}\right)^{n-1} \exp\left(-\frac{1}{G \sqrt{A^2 \cos^2 \theta + 4}}\right). \end{aligned} \quad (38)$$

Итак

$$\begin{aligned} \exp \frac{\bar{z}^2}{2} &= \exp \frac{1}{2} nG^2 \left[1 - \frac{1}{G \sqrt{A^2 \cos^2 \theta + 4}} + O\left(\frac{1}{n^2}\right) \right]^2 \sim \\ &\sim \exp\left(\frac{1}{2} nG^2 - \frac{2G}{\sqrt{A^2 \cos^2 \theta + 4}}\right) = \\ &= \exp\left[\frac{1}{2} n(1 + AG \cos \theta) - \frac{G}{\sqrt{A^2 \cos^2 \theta + 4}}\right], \end{aligned} \quad (39)$$

так как, возводя G в квадрат, находим, что $G^2 = 1 + AG \cos \theta$. Группируя слагаемые, получим

$$\begin{aligned} K &\sim T \sqrt{2\pi} \left(\frac{\sqrt{nG}}{e}\right)^{n-1} e^{n/2} \exp\left(-\frac{1}{G \sqrt{A^2 \cos^2 \theta + 4}}\right) - \\ &\quad - \frac{G}{\sqrt{A^2 \cos^2 \theta + 4}} - \frac{A^2 n}{2} \cos^2 \theta + \frac{n}{2} AG \cos \theta = \\ &= T \sqrt{2\pi} n^{(n-1)/2} G^{n-1} e^{-n/2} \exp\left(-\frac{n}{2} A^2 \cos^2 \theta + \frac{n}{2} AG \cos \theta\right), \end{aligned} \quad (40)$$

так как с помощью простых алгебраических преобразований можно показать, что слагаемые

$$1 - \frac{1}{G \sqrt{A^2 \cos^2 \theta + 4}} - \frac{G}{\sqrt{A^2 \cos^2 \theta + 4}}$$

в экспоненте взаимно уничтожаются. Коэффициент перед интегралом в формуле (34) при использовании асимптотического выражения для $\Gamma[(n+1)/2]$ асимптотически равен

$$\frac{(n-1) e^{-(\sin^2 \theta) (A^2 n)/2} \sin \theta^{n-2} e^{(n+1)/2}}{2^{n/2} \sqrt{\pi} \left(\frac{n+1}{2}\right)^{n/2} \sqrt{2\pi}}. \quad (41)$$

Комбинируя этот результат с предыдущими и группируя слагаемые (легко найти, что $T = G/\sqrt{1+G^2}$), имеем

$$-\frac{dQ}{d\theta} \sim \frac{n-1}{\sqrt{\pi n}} \frac{1}{\sqrt{1+G^2 \sin^2 \theta}} \left[G \sin \theta \exp \left(-\frac{A^2}{2} + \frac{1}{2} AG \cos \theta \right) \right]^n. \quad (42)$$

Это и есть искомое асимптотическое выражение для плотности вероятности $dQ/d\theta$.

Как установлено, коэффициент растет в основном как $\sqrt{n}$, однако имеется еще член вида $e^{-E_L(\theta)n}$, где

$$E_L(\theta) = A^2/2 - \frac{1}{2} AG \cos \theta - \log(G \sin \theta).$$

Можно показать, что если использовать для θ специальное значение $\theta_0 = \arccos A$ (см. рис. 4), то $E_L(\theta_0) = 0$, а также $E'_L(\theta_0) = 0$. Таким образом, для этого значения

$$\begin{aligned} G(\theta_0) &= \frac{1}{2} (A \cos \theta_0 + \sqrt{A^2 \cos^2 \theta_0 + 4}) = \\ &= \frac{1}{2} \left(\frac{A^2}{\sqrt{A^2+1}} + \sqrt{\frac{A^4}{A^2+1} + 4} \right) = \\ &= \frac{1}{2} \left(\frac{A^2}{\sqrt{A^2+1}} + \frac{A^2+2}{\sqrt{A^2+1}} \right) = \operatorname{cosec} \theta_0, \end{aligned}$$

и два слагаемых в логарифме взаимно уничтожаются. Далее

$$\frac{A^2}{2} - \frac{1}{2} AG \cos \theta_0 = \frac{A^2}{2} - \frac{1}{2} A \sqrt{A^2+1} \frac{A}{\sqrt{A^2+1}} = 0, \quad E_L(\theta_0) = 0.$$

Имеем также

$$E'_L(\theta) = \frac{1}{2} AG \sin \theta - \frac{1}{2} AG' \cos \theta - \frac{G'}{G} - \operatorname{ctg} \theta. \quad (43)$$

После значительных алгебраических преобразований слагаемое $-G'/G$ упрощается и оказывается равным

$$\frac{A \sin \theta}{\sqrt{A^2 \cos^2 \theta + 4}}.$$

Подставляя его и другие слагаемые, получим

$$\begin{aligned} E'_L(\theta) &= \frac{A^2}{2} \sin \theta \cos \theta + \frac{A^3 \cos^2 \theta \sin \theta}{4 \sqrt{A^2 \cos^2 \theta + 4}} + \\ &+ \frac{A}{4} \frac{(A^2 \cos^2 \theta + 4)}{\sqrt{A^2 \cos^2 \theta + 4}} \sin \theta + \frac{A \sin \theta}{\sqrt{A^2 \cos^2 \theta + 4}} - \operatorname{ctg} \theta. \end{aligned} \quad (44)$$

Суммируя и группируя слагаемые, упростим выражение, получив

$$\begin{aligned} E'_L(\theta) &= \frac{A}{2} (A \cos \theta + \sqrt{A^2 \cos^2 \theta + 4}) \sin \theta - \operatorname{ctg} \theta = AG \sin \theta - \operatorname{ctg} \theta = \\ &= \operatorname{ctg} \theta \left[\frac{A^2}{2} \sin^2 \theta + \frac{A}{2} \sin^2 \theta \sqrt{A^2 + \frac{4}{\cos^2 \theta} - 1} \right]. \end{aligned} \quad (45)$$

Заметим, что выражение в скобках является монотонно возрастающей функцией θ ($0 \leq \theta \leq \pi/2$), принимающей значение -1 при $\theta = 0$ и ∞ при $\theta = \pi/2$. Согласно ранее сказанному, $G = \operatorname{cosec} \theta_0$, $A = \operatorname{ctg} \theta_0$, так что $E'_L(\theta_0) = 0$; отсюда следует, что $E'_L(\theta) < 0$ для $0 \leq \theta < \theta_0$ и $E'_L(\theta) > 0$ для $\theta_0 \leq \theta < \pi/2$.

Из этого вытекает, что в области от θ_1 до $\pi/2$, когда $\theta_1 > \theta_0$, минимум $E_L(\theta)$ достигается при наименьшем в этой области значении θ , т. е. в θ_1 . Экспоненциальный член, входящий в оценку $Q(\theta)$, а именно $e^{-E_L(\theta)n}$, имеет для этой области максимум в точке θ_1 . В самом деле, для достаточно больших n максимум вышеуказанного выражения (45) должен быть вблизи θ_1 , поскольку присутствие n в экспоненте оказывает преобладающее влияние по сравнению с коэффициентом. Если обозначить этот коэффициент через $\alpha(\theta)$ и положить

$$y(\theta) = \alpha(\theta) e^{-nE_L(\theta)},$$

то

$$y'(\theta) = e^{-nE_L(\theta)} [-\alpha(\theta) nE'_L(\theta) + \alpha'(\theta)], \quad (46)$$

и, так как $\alpha(\theta) > 0$, то при достаточно большом n $y'(\theta)$ отрицательно и единственный максимум находится в θ_1 . В окрестности θ_1 функция экспоненциально убывает.

Теперь можно найти асимптотическую формулу для интеграла

$$Q(\theta) = \int_{\theta_1}^{\pi/2} \alpha(\theta) e^{-nE_L(\theta)} d\theta + Q(\pi/2). \quad (47)$$

Разбив интеграл на две части, получаем

$$Q(\theta) = \int_{\theta_1}^{\theta_1+n^{-2/3}} + \int_{\theta_1+n^{-2/3}}^{\pi/2} + Q(\pi/2). \quad (48)$$

На интервале определения первого интеграла $(1-\varepsilon)\alpha(\theta_1) \leq \alpha(\theta) \leq \alpha(\theta_1)(1+\varepsilon)$, и ε можно сделать сколь угодно малым, выбирая n достаточно большим. Это является следствием непрерывности и не обращения в нуль $\alpha(\theta)$ на указанном интервале. Используя разложение в ряд Тейлора с остаточным членом, имеем

$$e^{-nE_L(\theta)} = \exp \left[-nE_L(\theta_1) - n(\theta - \theta_1)E'_L(\theta_1) - n \frac{(\theta - \theta_1)^2}{2} E''_L(\theta^*) \right], \quad (49)$$

где θ^* находится в интервале между θ_1 и θ . При возрастании n максимальное значение остаточного члена, ограниченное величиной

$$n(n/2)^{-4/3} E''_{\max},$$

стремится к нулю. Следовательно, первый интеграл асимптотически совпадает с

$$\begin{aligned} & \alpha(\theta_1) \int_{\theta_1}^{\theta_1 + n^{-2/3}} \exp[-nE_L(\theta_1) - n(\theta - \theta_1)E'_L(\theta_1)] d\theta = \\ & = -\alpha(\theta_1) \exp[-nE_L(\theta_1)] \frac{\exp[-n(\theta - \theta_1)E'_L(\theta_1)]}{nE'_L(\theta_1)} \Big|_{\theta_1}^{\theta_1 + n^{-2/3}} \sim \\ & \sim \frac{\alpha(\theta_1) e^{-nE_L(\theta_1)}}{nE'_L(\theta_1)}, \end{aligned} \quad (50)$$

так как при больших n слагаемое, соответствующее верхнему пределу, сравнительно мало. Второй интеграл в пределах от $\theta_1 + n^{-2/3}$ до $\pi/2$ можно мажорировать значением подынтегральной функции в точке $\theta_1 + n^{-2/3}$, умноженной на интервал интегрирования

$$\frac{\pi}{2} - (\theta_1 + n^{-2/3})$$

(так как подынтегральная функция является монотонно убывающей при больших n). Значение подынтегральной функции в точке $\theta_1 + n^{-2/3}$ в соответствии с развитыми соображениями асимптотически стремится к величине

$$\alpha(\theta_1) \exp[-nE_L(\theta_1) - n(n^{-2/3})E'_L(\theta_1)],$$

малой по сравнению с первым интегралом [так как $Q(\pi/2) = \Phi(-A)$ в (47)]. Подставляя выражение для $\alpha(\theta_1)$ и заменяя θ_1 на θ , получаем асимптотическое выражение для

$$Q(\theta) \sim \frac{1}{\sqrt{n\pi}} \cdot \frac{\left[G \sin \theta \exp\left(-\frac{A^2}{2} + \frac{1}{2} AG \cos \theta\right) \right]^n}{\sqrt{1 + G^2 \sin \theta (AG \sin^2 \theta - \cos \theta)}} \quad (\pi/2 \geq \theta > \theta_0 = \text{arcctg } A). \quad (51)$$

Это выражение дает асимптотическую нижнюю границу для $P_{\text{еорт}}$, получаемую вычислением $Q(\theta)$ для θ_1 , такого, что

$$M\Omega(\theta_1) = \Omega(\pi).$$

Заметим, что асимптотическое выражение (51) можно преобразовать в асимптотическую формулу для функции нецентрального t -распределения подстановкой

$$\theta = \text{arcctg}(t/\sqrt{f}) \text{ и } n-1=f.$$

Это может быть полезным в других случаях применения нецентрального t -распределения.

7. Асимптотическое выражение для верхней границы при случайном кодировании

Найдем простое асимптотическое выражение для верхней границы $P_{e \text{ опт}}$ из формулы (20), основанное на методе случайного кодирования. Подставляя асимптотические выражения для $dQ(\theta)$ и $\Omega(\theta)/\Omega(\pi)$ в асимптотическую формулу для верхней границы, получим

$$Q(\theta_1) + e^{nR} \int_0^{\theta_1} \frac{\Gamma\left(\frac{n}{2}+1\right) (\sin \theta)^{n-1}}{n \Gamma\left(\frac{n+1}{2}\right) \pi^{1/2} \cos \theta} \times \\ \times \sqrt{\frac{n}{\pi}} \frac{\left[G \sin \theta \exp\left(-\frac{P}{2N} + \frac{1}{2} \sqrt{\frac{P}{N}} G \cos \theta\right) \right]^n}{\sqrt{1+G^2} \sin^2 \theta} d\theta. \quad (52)$$

Теперь требуется оценить интеграл

$$W = \int_0^{\theta_1} \frac{1}{\cos \theta \sin^3 \theta \sqrt{1+G^2}} \times \\ \times \exp\left\{n\left(-\frac{P}{2N} + \frac{1}{2} \sqrt{\frac{P}{N}} G \cos \theta + \log G + 2 \log \sin \theta\right)\right\} d\theta. \quad (53)$$

Положение здесь аналогично положению, возникающему при оценке $Q(\theta)$. Пусть коэффициент при n в экспоненте равен D . Заметив, что $D = -E_L(\theta) + \log \sin \theta$, находим для его производной

$$\frac{dD}{d\theta} = -AG \sin \theta + 2 \operatorname{ctg} \theta. \quad (54)$$

Уравнение $dD/d\theta = 0$ имеет единственный корень θ_c при $0 \leq \theta_c \leq \pi/2$ для любого фиксированного $A > 0$. Это следует из тех же соображений, которые использовались в связи с формулой (45) с той лишь разницей, что здесь имеется множитель 2 в слагаемом в правой части. Так, для $\theta < \theta_c$, производная $dD/d\theta$ является положительной, а D — возрастающей функцией θ . После достижения максимума D оказывается убывающей функцией.

Теперь можно решать задачу оценки интеграла (53) для различных случаев, характеризующихся различными соотношениями между θ_c и θ_1 .

Первый случай: $\theta_1 < \theta_c$. Здесь максимум экспоненты внутри области интегрирования получается при θ_1 . Следовательно, если n достаточно велико, максимум подинтегрального выражения имеет место в θ_1 . Асимптотическое значение можно оценить точно так же, как оценивалось $Q(\theta)$ в подобном случае. Интеграл разбивается на две части: первую от $\theta_1 - n^{-1/3}$ до θ_1 и вторую от 0 до $\theta_1 - n^{-1/3}$.

В первой части асимптотическое значение подинтегрального выражения имеет вид

$$\frac{1}{\cos \theta_1 \sin^3 \theta_1 \sqrt{1+G^2(\theta_1)}} \exp \left\{ n \left[-\frac{P}{2N} + \frac{1}{2} \sqrt{\frac{P}{N}} G(\theta_1) \cos \theta_1 + \log G(\theta_1) + 2 \log \sin \theta_1 - (\theta - \theta_1) [AG(\theta_1) \sin \theta_1 - 2 \operatorname{ctg} \theta_1] \right] \right\}. \quad (55)$$

Интеграл асимптотически равен

$$\frac{\exp \left\{ n \left[-\frac{P}{2N} + \frac{1}{2} \sqrt{\frac{P}{N}} G(\theta_1) \cos \theta_1 + \log G(\theta_1) + 2 \log \sin \theta_1 \right] \right\}}{\cos \theta_1 \sin^3 \theta_1 \sqrt{1+G^2(\theta_1)} [-AG(\theta_1) \sin \theta_1 + 2 \operatorname{ctg} \theta_1] n}. \quad (56)$$

Второй интеграл мал по сравнению с этим выражением, так как его можно мажорировать экспонентой с большим отрицательным показателем, умноженным на величину интервала $\theta_1 - n^{-2/3}$. Учитывая коэффициент

$$\frac{1}{\pi \sqrt{n}} \left[\frac{\Gamma\left(\frac{n}{2}+1\right)}{\Gamma\left(\frac{n+1}{2}\right)} \right] e^{nR}$$

и используя соотношение

$$\frac{\Gamma\left(\frac{n}{2}+1\right)}{\Gamma\left(\frac{n+1}{2}\right)} \sim \sqrt{\frac{n}{2}},$$

видим, что главный член приближенно равен

$$\frac{\left[G \sin \theta_1 \exp \left(-\frac{A^2}{2} + \frac{1}{2} AG \cos \theta_1 \right) \right]^n}{\sqrt{n\pi} \sqrt{1+G^2} \sin \theta_1 (2 \cos \theta_1 - AG \sin^2 \theta_1)}. \quad (57)$$

Сопоставляя этот результат с ранее полученным асимптотическим выражением (51), получаем для $Q(\theta_1)$ следующее асимптотическое выражение для верхней границы $P_{e \text{ opt}}$ при $\theta_1 < \theta_c$:

$$\left(1 - \frac{\cos \theta_1 - AG \sin^2 \theta_1}{2 \cos \theta_1 - AG \sin^2 \theta_1} \right) \frac{\left[G \sin \theta_1 \exp \left(-\frac{A^2}{2} + \frac{1}{2} AG \cos \theta_1 \right) \right]^n}{\sqrt{n\pi} \sqrt{1+G^2} \sin \theta_1 (AG \sin^2 \theta_1 - \cos \theta_1)}. \quad (58)$$

Для нашей нижней границы асимптотически справедливо то же самое выражение, но без множителя в скобках. Таким образом, оба эти асимптотических выражения разнятся лишь на множитель

$$\left(1 - \frac{\cos \theta_1 - AG \sin^2 \theta_1}{2 \cos \theta_1 - AG \sin^2 \theta_1} \right),$$

не зависящий от n . Этот множитель увеличивается, когда θ_1 возрастает от величины θ_0 , соответствующей пропускной способности канала, до критического значения θ_c , при котором знаменатель обращается

в нуль. В этих границах множитель возрастает от единицы до бесконечности. Другими словами, для больших n вероятность P_{opt} определяется с точностью до множителя. Более того, для скоростей передачи, близких к пропускной способности канала, величина этой неопределенности стремится к нулю по мере стремления скорости передачи к значению пропускной способности канала. Очень интересно, что эти кажущиеся на первый взгляд слабыми оценки для некоторых областей значений переменных дают такую точную информацию.

Второй случай: $\theta_1 > \theta_c$. Для θ_1 в этом случае предыдущие рассуждения не пригодны, так как максимум экспоненты не находится на конце интервала интегрирования, а расположен внутри него. Этот единственный максимум получается при значении θ — корне уравнения $2 \cos \theta_c - AG \sin^2 \theta_c = 0$. Разобьем область интегрирования на три части: от 0 до $\theta_c - n^{-2/3}$, от $\theta_c - n^{-2/3}$ до $\theta_c + n^{-2/3}$ и от $\theta_c + n^{-2/3}$ до θ . Используя весьма простые соображения, видим, что в окрестности значения θ_c экспонента имеет вид

$$\exp \left(-n \left\{ E_L(\theta_c) + \frac{(\theta - \theta_c)^2}{2} E_L''(\theta_c) + O[(\theta - \theta_c)^3] \right\} \right).$$

Коэффициент перед экспонентой можно считать постоянным на малом интервале интегрирования вокруг θ_c . Интеграл (53) по этой части асимптотически равен

$$\frac{1}{\cos \theta_c \sin^3 \theta_c \sqrt{1+G^2}} \int \exp \left\{ -n \left[E_L(\theta_c) + \frac{(\theta - \theta_c)^2}{2} E_L''(\theta_c) \right] \right\} d\theta \sim \\ \sim \frac{1}{\cos \theta_c \sin^3 \theta_c \sqrt{1+G^2}} \exp[-n E_L(\theta_c)] \frac{\sqrt{2\pi}}{\sqrt{n E_L''(\theta_c)}}. \quad (59)$$

Другие два интеграла малы при больших n по сравнению с вычисленным. Это видно из соображений, близких к вышеприведенным. Они оцениваются величиной подынтегральной функции на конце области вблизи θ_c , помноженной на величину интервала интегрирования. В целом интеграл в (52) имеет следующее асимптотическое значение:

$$\frac{1}{\sqrt{\pi n} \cos \theta_c \sin^3 \theta_c \sqrt{1+G^2} \sqrt{E_L''(\theta_c)}} e^{-n[E_L(\theta_c) - R]}. \quad (60)$$

Другое слагаемое в интеграле (52), а именно $Q(\theta_1)$, асимптотически мало по сравнению с только что полученным для случая $\theta > \theta_c$, так как коэффициент при n в экспоненте выражения $Q(\theta)$ в уравнении (51) меньше. Все это приводит к тому, что граница, оцененная методом случайного кодирования, асимптотически равна

$$\frac{1}{\cos \theta_c \sin^3 \theta_c \sqrt{\pi n E_L''(\theta_c) [1+G(\theta_c)^2]}} e^{-n[E_L(\theta_c) - R]} \quad (61)$$

для $\theta > \theta_c$ или для скорости передачи R , меньшей R_c , скорости передачи, соответствующей θ_c . Заметим, что скорость передачи R_c почти на полбита меньше пропускной способности канала, если $A \gg 4$, и стремится к точному совпадению с ней при $A \rightarrow \infty$. Для меньших значений A разность $C - R_c$ уменьшается, но отношение C/R_c стремится к 4, когда $A \rightarrow 0$.

8. Жесткая верхняя граница $P_{e \text{ opt}}$

В этом разделе путем преобразования формулы (20) будут найдены формулы для верхней границы вероятности ошибки, справедливые при всех n . Определим сначала верхнюю границу для $Q'(\theta)$. В работе Девида и Крускала¹⁾ интеграл (35) преобразуется в умноженный на $\bar{z}^v \exp(-\frac{1}{2} \bar{z}^2 + \bar{z} \omega \sqrt{v+1})$ интеграл (в обозначениях указанной работы)

$$U = \int_{-\infty}^{\infty} \varphi_{\bar{z}}(y) \exp \left\{ -\frac{1}{2} y^2 + v \left[\ln \left(1 + \frac{y}{\bar{z}} \right) - \frac{y}{\bar{z}} \right] \right\} dy.$$

Отметим, что подинтегральное выражение может быть оценено сверху значением $e^{-y^2/2}$. Это выясняется в абзаце указанной работы, содержащем равенство (26). Поэтому этот интеграл может быть оценен сверху через $\sqrt{2\pi}$, а наш интеграл в выражении (34), содержащийся в $dQ/d\theta$, оценивается сверху следующим выражением:

$$\begin{aligned} & \int_0^{\infty} \exp \left[-\frac{(r - A \sqrt{n} \cos \theta)^2}{2} \right] r^{n-1} dr = \\ & = \left(\frac{\bar{z}}{e} \right)^{n-1} \exp \left(\frac{\bar{z}}{2} \right)^2 \exp \frac{A^2 n}{2} \cos^2 \theta U \leq \\ & \leq \left(\frac{\bar{z}}{e} \right)^{n-1} \exp \left(\frac{\bar{z}}{2} \right)^2 \exp \frac{A^2 n}{2} \cos^2 \theta \sqrt{2\pi}. \end{aligned}$$

Имеем

$$\bar{z} = \frac{1}{2} \sqrt{n} (A \cos \theta + \sqrt{A^2 \cos^2 \theta + 4 - 4/n}) \leq \sqrt{n} G.$$

Подстановка вместо $\bar{z}$ этой большей величины дает

$$\left(\frac{\sqrt{n} G}{e} \right) \exp \left(\frac{n G^2}{2} - \frac{A^2 n}{2} \cos^2 \theta \right) \sqrt{2\pi}.$$

¹⁾ См. сноску на стр. 558.

Теперь имеем

$$-\frac{dQ}{d\theta} \leq \frac{(n-1) \exp\left(\frac{-A^2 n}{2} \sin^2 \theta\right) (\sin \theta)^{n-2}}{2^{n/2} \sqrt{\pi} \Gamma\left(\frac{n+1}{2}\right)} \times \\ \times \left(\frac{\sqrt{n} G}{e}\right)^{n-1} \exp\left(\frac{n G^2}{2} - \frac{A^2 n}{2} \cos^2 \theta\right) \sqrt{2\pi}. \quad (62)$$

Заменяя Γ -функцию ее стирлинговским приближением

$$\left(\frac{n+1}{2}\right)^{n/2} \exp\left(\frac{n+1}{2}\right) \sqrt{2\pi}$$

(которое всегда меньше ее), а также подставляя $\sqrt{2}$ вместо $[1 + (1/n)]^{n/2}$ (где снова второе меньше первого), лишь увеличим правую часть неравенства. Поэтому

$$-\frac{dQ}{d\theta} \leq \frac{(n-1) (G \sin \theta)^n \exp\left[\left(\frac{n}{2}\right) (-A^2 + 1 + AG \cos \theta)\right]}{\sqrt{n} G \sin^2 \theta \sqrt{2\pi} \exp\left(\frac{n-3}{2}\right)} \leq \\ \leq \frac{(n-1) e^{3/2} e^{-E_L(\theta)n}}{\sqrt{2\pi n} G \sin^2 \theta}. \quad (63)$$

Заметим, что последняя оценка отличается от асимптотического выражения (42) лишь множителем

$$\frac{e^{3/2} \sqrt{1+G^2}}{\sqrt{2} G} \leq e^{3/2}$$

(так как $G \geq 1$). Теперь может быть установлена жесткая верхняя граница для $Q(\theta)$

$$Q(\theta_1) = \int_{\theta_1}^{\pi/2} -\frac{dQ}{d\theta} d\theta + Q\left(\frac{\pi}{2}\right).$$

Используем найденное θ_1 для выражения верхней границы $dQ/d\theta$ в интеграле. Коэффициент при $-n$ в показателе e

$$E_L(\theta) = \frac{1}{2} (A^2 - AG \cos \theta) - \lg G \sin \theta$$

положителен и монотонно возрастает с ростом θ при $\theta > \theta_0$, как было показано ранее. Его производная равна

$$E'_L(\theta) = AG \sin \theta - \operatorname{ctg} \theta.$$

На рис. 5 эта производная как функция θ изображена в виде кривых, которые либо монотонно возрастают от $-\infty$ при $\theta=0$

до A при $\theta = \pi/2$, либо имеют единственный максимум. Во всех случаях кривые выпуклы вверх. Чтобы показать это аналитически, возьмем вторую производную от E'_L . Она состоит из суммы отрицательных слагаемых.

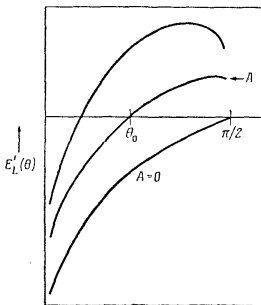
Возвращаясь к нашей верхней границе Q , видим, что коэффициент в выражении (63) не превышает

$$\frac{\frac{\sqrt{n}}{\sqrt{\pi}} e^{3/2}}{\sin^2 \theta_1},$$

где $\sin \theta$ и G заменены на $\sin \theta_1$ и 1 — минимальные значения этих величин в рассматриваемых областях. Заменим теперь $e^{-nE_L(\theta)}$ на

$$\exp\{-n[E_L(\theta_1) + (\theta - \theta_1)h]\}.$$

Если h выбрано равным минимуму $E'_L(\theta)$, то такая замена приводит к возрастанию интеграла и, следовательно, даст верхнюю границу.



Р и с. 5. $E_L(\theta)$ как функция θ .

Из поведения $E'_L(\theta)$ видно, что этот минимум имеет место либо при θ_1 , либо при $\pi/2$. Поэтому имеем $h = \min\{A, AG(\theta_1) \sin \theta_1 - \text{ctg } \theta_1\}$. При такой замене интеграл превращается в простой экспоненциальный и его можно непосредственно вычислить.

Член $Q(\pi/2)$ имеет вид

$$\Phi(-A\sqrt{n}) \leq \frac{1}{\sqrt{2\pi n} A} e^{-A^2 n/2}.$$

Если вести интегрирование до бесконечности, вместо того чтобы остановиться на $\pi/2$, то лишняя часть даст добавку, с избытком компенсирующую $Q(\pi/2)$. Действительно, $E_L(\pi/2) = A^2/2$, так что лишняя добавка составляет по меньшей мере

$$\frac{\sqrt{n} e^{3/2}}{An \sin^2 \theta_1 \sqrt{2\pi}} e^{-A^2 n/2},$$

если интегрировать

$$\frac{\sqrt{n} e^{3/2}}{\sin^2 \theta_1 \sqrt{2\pi}} e^{-A^2 n/2 - n(\theta - \theta_1)A}$$

до бесконечности, вместо того чтобы остановиться на $\pi/2$. Поскольку $e^{3/2}/\sin^2 \theta_1 \geq 1$, можно отбросить член $Q(\pi/2)$, компенсируя этим лишнюю добавку в интеграле.

Следовательно, можно ограничить $Q(\theta_1)$ следующим образом:

$$Q(\theta_1) \leq \frac{e^{3/2} \exp \{ (n/2) [AG(\theta_1) \cos \theta_1 - A^2 + 2 \log G \sin \theta_1] \}}{\sqrt{2\pi n} \sin^2 \theta_1 \min(A, AG(\theta_1) \sin \theta_1 - \operatorname{ctg} \theta_1)}. \quad (64)$$

Чтобы ограничить сверху $P_{e \text{ opt}}$, согласно выражению (3), необходимо найти верхнюю границу слагаемого

$$\int_0^{\theta_1} \frac{\Omega(\theta)}{\Omega(\theta_1)} dQ(\theta).$$

Это можно сделать методом, очень сходным с тем, который использовался для оценки $\int dQ(\theta)$. Сначала ограничим $\Omega(\theta)/\Omega(\theta_1)$ сверху, используя выражение (21). Имеем

$$\begin{aligned} \frac{\Omega(\theta)}{\Omega(\theta_1)} &= \frac{\int_0^{\theta} (\sin x)^{n-2} dx}{\int_0^{\theta_1} (\sin x)^{n-2} dx} = \frac{\int_0^{\theta} (\sin x)^{n-2} dx}{\int_0^{\theta} (\sin x)^{n-2} dx + \int_{\theta}^{\theta_1} (\sin x)^{n-2} dx} \ll \\ &\ll \frac{\int_0^{\theta} (\sin x)^{n-2} \cos x dx}{\int_0^{\theta} (\sin x)^{n-2} \cos x dx + \cos \theta \int_{\theta}^{\theta_1} (\sin x)^{n-2} dx} \ll \\ &\ll \frac{\int_0^{\theta} (\sin x)^{n-2} \cos x dx}{\int_0^{\theta} (\sin x)^{n-2} \cos x dx + \int_{\theta}^{\theta_1} (\sin x)^{n-2} \cos x dx} \end{aligned}$$

и окончательно

$$\frac{\Omega(\theta)}{\Omega(\theta_1)} \leq \frac{(\sin \theta)^{n-1}}{(\sin \theta_1)^{n-1}}. \quad (65)$$

Здесь неравенство во второй строке получается из-за того, что первый интеграл в знаменателе второй строки преобразуется так же, как и числитель, а второй интеграл в знаменателе уменьшается сильнее, так как $\cos \theta$ — убывающая функция. В последней строке знаменатель уменьшается еще больше из-за введения косинуса под знак интеграла.

Используя это неравенство, а также формулу верхней границы (63) для $dQ/d\theta$, имеем

$$\begin{aligned} \int_0^{\theta_1} \frac{\Omega(\theta)}{\Omega(\theta_1)} dQ(\theta) &\leq \int_0^{\theta_1} \frac{(\sin \theta)^{n-1}}{(\sin \theta_1)^{n-1}} \frac{(n-1) e^{3/2} (G \sin \theta)^n e^{(n/2)(-A^2 + A \cos \theta G)}}{\sqrt{2\pi n} G \sin^2 \theta} d\theta = \\ &= \frac{(n-1) e^{3/2}}{\sqrt{2\pi n} (\sin \theta_1)^{n-1}} \int_0^{\theta_1} G^n (\sin \theta) 2n^{-3} e^{(n/2)(-A^2 + A \cos \theta G)} d\theta. \end{aligned} \quad (66)$$

Вблизи точки θ_1 подинтегральное выражение при больших n ведет себя как экспонента (полагаем $\theta_1 < \theta_c$), и поэтому имеет смысл искать жесткую верхнюю границу в форме

$$\frac{k}{\sqrt{n}} e^{-E_L(\theta_1)n},$$

где k не зависит от n . Это ведет, однако, к значительным трудностям, ввиду чего ограничимся следующими грубыми построениями.

Подинтегральная функция ограничена своим максимальным значением. Если $\theta_1 < \theta_c$, то максимум подинтегрального выражения достигается по крайней мере при достаточно большом n в точке θ_1 . В этом случае интеграл ограничен величиной

$$\theta_1 G^n(\theta_1) (\sin \theta_1)^{2n-3} e^{(n/2)[-A^2 + A \cos \theta_1 G(\theta_1)]}.$$

Полное выражение для $P_{e \text{ opt}}$ может быть ограничено выражением (67) [оно получается прибавлением найденной оценки к границе (64) для $Q(\theta_1)$]

$$P_{e \text{ opt}} \leq \frac{\sqrt{n} e^{3/2} \theta_1 e^{-E_L(\theta_1)}}{\sqrt{2\pi} \sin^2 \theta_1} \left\{ 1 + \frac{1}{n \theta_1 \min[A, AG(\theta_1) \sin \theta_1 - \text{ctg} \theta_1]} \right\}. \quad (67)$$

Необходимо помнить, что выражение (67) справедливо лишь в том случае, если $\theta_1 < \theta_c$ и n достаточно велико для того, чтобы максимум рассмотренного выше подинтегрального выражения получался при θ_1 . Можно также определить границы для $\theta_1 > \theta_c$, опираясь на максимальное значение подинтегрального выражения.

9. Жесткая нижняя граница для $P_{e \text{ opt}}$

В этом разделе будет найдена нижняя граница для $P_{e \text{ opt}}$, имеющая место при всех n . Для этого определим сначала нижнюю границу для $Q'(\theta)$, а по ней найдем нижнюю границу для $Q(\theta)$. Метод ее нахождения очень сходен с методом, примененным для отыскания верхней границы.

В работе Девида и Крускала¹⁾ нахождение приведенного выше интеграла (35) сводится к вычислению следующего интеграла [выражение (2.5) указанной работы]:

$$\begin{aligned} \int_{-\frac{v}{z}}^{\infty} \left(1 + \frac{y}{z}\right)^v \exp\left(-\frac{1}{2}y^2 - y\frac{v}{z}\right) dy &\geq \\ &\geq \int_0^{\infty} \exp\left\{-\frac{1}{2}y^2 + v\left[\ln\left(1 + \frac{y}{z}\right) - \frac{y}{z}\right]\right\} dy \geq \\ &\geq \int_0^{\infty} \exp\left[-\frac{1}{2}y^2 + v\left(\frac{-y^2}{2z^2}\right)\right] dy = \\ &= \int_0^{\infty} \exp\left[\frac{-y^2}{2}\left(1 + \frac{v}{z^2}\right)\right] dy = \frac{1}{2} \frac{\sqrt{2\pi}}{\sqrt{1 + \frac{v}{z^2}}} \geq \frac{\sqrt{2\pi}}{2\sqrt{2}} = \frac{\sqrt{\pi}}{2}. \end{aligned}$$

Здесь использовалось неравенство

$$\ln\left(1 + \frac{y}{z}\right) - \frac{y}{z} \geq -\frac{y^2}{2z^2} \quad \text{для } \frac{y}{z} > 0,$$

а также то обстоятельство, что $v/z^2 \leq 1$. Последнее следует из соотношения (2.3) указанной работы, деленного на $\bar{z}^2$.

Используя эту нижнюю границу, получим из выражения (34)

$$\frac{dQ}{d\theta} \geq \frac{(n-1) \sin^{n-2}\theta \exp\left(\frac{-A^2 n}{2}\right)}{2^{n/2} \sqrt{\pi} \Gamma\left(\frac{n+1}{2}\right)} \left(\frac{\bar{z}}{e}\right)^{n-1} \exp\left(\frac{\bar{z}^2}{2}\right) \frac{\sqrt{\pi}}{2}. \quad (68)$$

Замечая, что $z \geq \sqrt{n-1} G$ и что Γ

$$\left(\frac{n+1}{2}\right) < \left(\frac{n+1}{2}\right)^{n/2} e^{-(n+1)/2} \sqrt{2\pi} \exp\left[\frac{1}{6(n+1)}\right],$$

и используя неравенство

$$\left(\frac{n-1}{n+2}\right)^{n/2} \geq \frac{1}{3} \quad \text{для } n \geq 2,$$

¹⁾ David H. T., Kruskal W. H., The WAGR sequential t -test reaches a decision with probability one, *Ann. Math. Stat.*, 27, September (1956), 797.

получим

$$\frac{dQ}{d\theta} \geq \frac{1}{6\sqrt{2\pi}} \left\{ \frac{\sqrt{n-1} e^{3/2} e^{-nE_L(\theta)}}{G \exp\left[\frac{G^2}{2} + \frac{1}{6(n+1)}\right] \sin^2\theta} \right\} \quad \text{для } n \geq 2. \quad (69)$$

Это и есть нижняя граница $dQ/d\theta$.

Чтобы получить нижнюю границу $Q(\theta)$, можно использовать тот же самый прием, что и выше. Здесь, однако, коэффициенты заменяются их минимальными значениями в рассматриваемой области, а экспонента — значением $-nE_L(\theta_1) - n(\theta - \theta_1)E'_L \max$

$$\begin{aligned} E'_L &= AG \sin \theta - \operatorname{ctg} \theta \leq \\ &\leq AG \leq \\ &\leq A(A+1). \end{aligned}$$

Аналогично значение G ограничивается сверху величиной $(A+1)$, а $\sin^2 \theta$ — единицей. Таким образом, получаем

$$Q(\theta_1) \geq \int_{\theta_1}^{\pi/2} \frac{\sqrt{n-1} e^{3/2} e^{-nE_L(\theta_1)} e^{-n(\theta - \theta_1)A(A+1)}}{6\sqrt{2\pi} (A+1) \exp\left[\frac{(A+1)^2}{2} + \frac{1}{6(n+1)}\right]} d\theta + Q\left(\frac{\pi}{2}\right). \quad (70)$$

Интегрируя и замечая, что член с подстановкой предела $\pi/2$ может быть включен в $Q(\pi/2) - \operatorname{erf} A$, приходим к выражению нижней границы:

$$Q(\theta_1) \geq \frac{\sqrt{n-1} e^{3/2} e^{-nE_L(\theta_1)}}{6\sqrt{2\pi} (A+1)^3 \exp\left[\frac{(A+1)^2}{2} + \frac{1}{6(n+1)}\right]}. \quad (71)$$

10. Поведение границ вблизи пропускной способности канала

Как мы видели, вблизи пропускной способности канала верхняя и нижняя асимптотические границы сближаются. Если нижнюю асимптотическую границу [см. соотношение (42)] представить в виде разложения Тейлора по θ вблизи θ_0 , ограничиваясь членом $(\theta - \theta_0)^2$, то получим выражение, которое можно использовать в окрестности пропускной способности канала. Другое приближение получим, возвращаясь к нецентральному t -распределению и используя его нормальную аппроксимацию, которая дает хорошее приближение вблизи среднего значения, как показано в работе¹⁾.

¹⁾ Johnson N. L., Welch B. L., Applications of the noncentral t -distribution, *Biometrika*, 31 (1939), 362.

Эти методы дают в указанной окрестности следующие аппроксимации [поскольку $E(\theta_0) = E'(\theta_0) = 0$]:

$$-\frac{dQ}{d\theta} \approx \frac{\sqrt{n}(1+A^2)}{\sqrt{\pi}\sqrt{2+A^2}} \exp\left[-n \frac{(A^2+1)^2}{A^2+2}(\theta-\theta_0)^2\right],$$

$$Q(\theta) \approx \Phi\left[(\theta_0-\theta) \frac{A^2+1}{\sqrt{A^2+2}} \sqrt{2n}\right] \quad (72)$$

или (так как вблизи пропускной способности канала, как было показано, $e^{-R} \approx x \sin \theta$)

$$\theta - \theta_0 \approx A^{-1}(C - R),$$

$$P_{e \text{ opt}}\left(n, R, \sqrt{\frac{P}{N}}\right) \approx \Phi\left[\sqrt{2n} A^{-1} \frac{A^2+1}{\sqrt{A^2+2}}(R-C)\right] =$$

$$= \Phi\left[\frac{P+N}{\sqrt{P(P+2N)}} \sqrt{2n}(R-C)\right]. \quad (73)$$

Кривые надежности вблизи значения C приближенно выражаются формулой

$$E_-(R) \approx \frac{(P+N)^2}{P(P+2N)}(C-R)^2. \quad (74)$$

Интересно, что Райс¹⁾ произвел оценку поведения величины E , стоящей в экспоненте нижней границы для случая близости скорости передачи к пропускной способности канала. В наших обозначениях он получил

$$E^*(R) \approx \frac{P+N}{2P}(C-R)^2,$$

т. е. меньшее значение, чем дает формула (74), так что у него получается большая длина кодов для достижения той же вероятности ошибки. Очевидно, что разница происходит из-за несколько различных способов построения случайных кодов. Райсовские коды образуются размещением точек согласно n -мерному гауссовскому распределению, каждая из координат которого имеет дисперсию P . В наших кодах точки размещаются случайно на сфере с точно фиксированным радиусом $\sqrt{nP}$. Это, конечно, почти одно и то же, если n достаточно велико, так как при райсовском подходе точки с вероятностью, стремящейся к единице, лежат между сферами радиусов $\sqrt{nP}(1-\epsilon)$ и $\sqrt{nP}(1+\epsilon)$ (при любом $\epsilon > 0$). Однако мы рассматриваем события, имеющие очень малую вероятность, всегда, когда оцениваем вероятности ошибок. Поэтому точки внутри сферы достаточно сильно влияют на экспоненту E . Другими словами, райсовский способ построения кодов достаточен, чтобы дать коды, обеспечивающие приближение вероятности ошибки к нулю при скоро-

¹⁾ Rice S. O., Communication in the presence of noise—probability of error for two encoding schemes, *BSTJ*, 29, January (1950), 60.

стях передачи, стремящихся к пропускной способности канала. Однако эта сходимость будет все же медленнее (даже в экспоненте), чем максимально достижимая. Для достижения наилучшего возможного E , очевидно, необходимо избегать наличия слишком большого числа кодовых точек внутри сферы радиуса $\sqrt{nP}$.

При скорости передачи R , большей, чем значение пропускной способности канала, имеем $\theta_1 < \theta_0$. Так как Q -распределение аппроксимируется нормальным распределением со средним θ_0 и дисперсией $2n(A^2 + 1)^2/(A^2 + 2)$, то $Q(\theta_1)$ стремится к 1 с ростом n при любой фиксированной скорости передачи, большей C . Далее, даже если скорость передачи R меняется, оставаясь всегда больше C (возможно, приближаясь к ней сверху при возрастании n), то все еще будет иметь место неравенство $P_{e \text{ opt}} > 1/2 - \varepsilon$ для $\varepsilon > 0$ и достаточно большого n .

11. Верхняя граница для $P_{e \text{ opt}}$, полученная методом исчерпывания

При малой скорости передачи, когда верхняя и нижняя границы расходятся, можно получить лучшие оценки другим методом. Можно показать, что при достаточно малой скорости передачи главный вклад в вероятность ошибки дают наиболее близкие друг к другу и потому часто перепутываемые друг с другом кодовые точки. Общая же средняя структура кода менее важна. Особо существенно при малой скорости передачи максимизировать минимальное расстояние между соседними точками. Определение верхней и нижней границ при малой скорости передачи будет основано на этих соображениях.

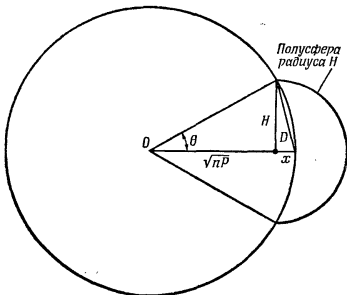
Сначала будет показано, что при $D \leq \sqrt{2nP}$ можно найти по меньшей мере

$$M_D = \left(\sin 2 \arcsin \frac{D}{2\sqrt{nP}} \right)^{1-n}$$

точек на поверхности n -мерной сферы радиуса $\sqrt{nP}$, таких, что ни одна пара из них не имеет между собой расстояния, меньшего D . (Если M_D не целое число, то берется ближайшее к нему наибольшее целое число.) Применяемый метод сходен с использованным ранее Гилбертом для случая двоичного симметричного канала.

Выберем некоторую точку на поверхности сферы как первую точку. Выбросим с поверхности сферы все точки с расстоянием, меньшим или равным D от выбранной точки. На рис. 6 буква x обозначает выбранную точку, а выброшенная площадь соответствует пересечению окружности с конусом. Непосредственно видно, что эта площадь, конечно, меньше (при $D \leq \sqrt{2nP}$) площади полусферы радиуса H

и тем более меньше площади всей сферы радиуса H . Если такое выбрасывание не исчерпывает начальную сферу, снова выбираем точку в оставшейся части и выбрасываем точки с расстоянием от этой новой точки, меньшим D . Это опять не удалит площади больше, чем площадь сферы радиуса H . Продолжаем действовать таким образом, пока не останется ни одной невыброшенной точки. Заметим, что



Р и с. 6. Геометрия сферы радиуса $\sqrt{nP}$.

каждая выбранная следующая точка удалена от предыдущей не менее чем на расстояние D . Следовательно, все расстояния между точками не меньше D . Далее замечаем, что число повторений этой операции может по крайней мере равняться отношению поверхности сферы радиуса $\sqrt{nP}$ к поверхности сферы радиуса H . Отношение площадей этих поверхностей, очевидно, равно

$$\left(\frac{\sqrt{nP}}{H} \right)^{n-1}.$$

Из простых геометрических построений (рис. 6) видно, что H и D связаны следующим образом:

$$\begin{aligned} \sin \theta &= \frac{H}{\sqrt{nP}}, \\ \sin \frac{\theta}{2} &= \frac{D}{2\sqrt{nP}}, \end{aligned}$$

следовательно,

$$H = \sqrt{nP} \sin 2 \arcsin \frac{D}{2\sqrt{nP}}. \quad (75)$$

Подставляя H в предыдущее отношение, видим, что можно разместить по меньшей мере

$$M_D = \left(\sin 2 \arcsin \frac{D}{2 \sqrt{nP}} \right)^{-(n-1)}$$

точек с расстояниями друг от друга, не меньшими D для $D \leq \sqrt{2nP}$.

Если имеется M_D точек с минимальным расстоянием не меньше D , то вероятность ошибки при оптимальном декодировании меньше или равна

$$M_D \Phi \left(\frac{-D}{2 \sqrt{N}} \right).$$

Чтобы показать это, просуммируем, предположив наихудший случай, вероятности для каждой из точек быть принятой за любую другую точку.

Вероятность для точки 1 оказаться ближе к точке 2, чем к первоначальной точке 1, равна $\Phi \left[-D/(2 \sqrt{N}) \right]$, т. е. вероятности того, что точка сместится не менее чем на $D/2$ (половина минимального расстояния). Вклад в ошибку по этой причине не может, следовательно, превосходить $(1/M_D) \Phi \left[-D/(2 \sqrt{N}) \right]$ (множитель $1/M_D$ соответствует вероятности передачи сообщения 1). Аналогичные соображения имеют место и для других упорядоченных пар точек, вносящих всего $M_D(M_D - 1)$ вкладов такого рода. Следовательно, вероятность ошибки не может превзойти $(M_D - 1) \Phi \left[-D/(2 \sqrt{N}) \right]$ или, упрощая, $M_D \Phi \left[-D/(2 \sqrt{N}) \right]$.

Если положить

$$e^{nR} = M_D = \left(\sin 2 \arcsin \frac{D}{2 \sqrt{nP}} \right)^{-(n-1)},$$

то скорость передачи (в натуральных единицах) выразится формулой

$$R = \left(1 - \frac{1}{n} \right) \log \left(\sin 2 \arcsin \frac{D}{2 \sqrt{nP}} \right) - 1$$

при

$$P_e \leq e^{nR} \Phi \left(\frac{-D}{2 \sqrt{N}} \right) \leq e^{nR} \frac{\sqrt{2N}}{D \sqrt{\pi}} e^{-(D^2/8N)}, \quad (76)$$

что следует из хорошо известной формулы для верхней границы

$$\Phi(-x) \leq \left(\frac{1}{x \sqrt{2\pi}} \right) e^{-x^2/2}.$$

Это — параметрические уравнения относительно D . Более удобно

положить $D = \lambda \sqrt{2nP}$. При этом будем иметь

$$R = \left(1 - \frac{1}{n}\right) \log \left(\sin 2 \arcsin \frac{\lambda}{\sqrt{2}}\right)^{-1},$$

$$P_e \leq \frac{1}{\lambda \sqrt{\frac{\pi n P}{n}}} e^{n[R - (\lambda^2 P)/(4N)]}. \quad (77)$$

Асимптотическая надежность, т. е. коэффициент при $-n$ в экспоненте P_e , равна $(\lambda^2 P/4N) - R$. При $n \rightarrow \infty$ она стремится к

$$\left(\sin \frac{1}{2} \arcsin e^R\right)^2 \frac{P}{2N} - R.$$

Асимптотическую нижнюю границу надежности находим путем исключения λ

$$E \geq \left(\sin \frac{1}{2} \arcsin e^R\right)^2 \frac{P}{2N} - R. \quad (78)$$

При $R \rightarrow 0$ выражение, стоящее справа, стремится к $P/(4N)$.

Эта нижняя граница экспоненты изображена на кривых в разд. 14. По ним можно видеть, что при малых скоростях передачи эта граница дает больше информации, чем граница, полученная методом случайного кодирования. Можно, однако, улучшить метод случайного кодирования так называемым процессом «вычеркивания». Он приводит к границе, совпадающей с только что найденной и даже несколько более сильной внутри части области. Не входя детально в описание этого метода, упомянем лишь, что он состоит в вычеркивании из случайного кода точек ансамбля, которые имеют слишком близкое соседство друг к другу, и в рассмотрении кодов, которые возникают после такого вычеркивания.

12. Нижняя граница P_e для гауссовского канала, получаемая исходя из минимального расстояния

Пусть m_{is} ($i = 1, 2, \dots, M$; $s = 1, 2, \dots, n$) есть s -я координата кодового слова в коде длины n с M кодовыми словами. Будем полагать, что средняя мощность ограничена P , так что

$$\frac{1}{nM} \sum_{i,s} m_{is}^2 \leq P. \quad (79)$$

Положим также, что независимый гауссовский шум мощности N действует аддитивным образом по всем координатам.

Вычислим средний квадрат расстояния между всеми $M(M-1)/2$ парами точек в n -мерном пространстве, соответствующих M кодовым

словам. Квадрат расстояния между словом i и словом j равен

$$\sum_s (m_{is} - m_{js})^2.$$

Среднее значение $\bar{D}^2$ по всем парам равно

$$\bar{D}^2 = \frac{1}{M(M-1)} \sum_{s, i, j} (m_{is} - m_{js})^2.$$

Заметим, что каждое расстояние учитывается дважды в сумме, а также, что члены, содержащиеся в сумме, при $i = j$ дают вклад, равный нулю. Возведем в квадрат члены суммы

$$\begin{aligned} \bar{D}^2 &= \frac{1}{M(M-1)} \left(\sum_{i, j, s} m_{is}^2 - 2 \sum_s \sum_{i, j} m_{is} m_{js} + \sum_{i, j, s} m_{js}^2 \right) = \\ &= \frac{1}{M(M-1)} \left[2M \sum_{i, s} m_{is}^2 - 2 \sum_s \left(\sum_i m_{is} \right)^2 \right] \leq \\ &\leq \frac{1}{M(M-1)} 2MPnM; \quad \bar{D}^2 \leq \frac{2nMP}{M-1}. \end{aligned} \quad (80)$$

Здесь третья строка получена использованием неравенства (79) для средней мощности, а также того обстоятельства, что второй член во второй строке необходимо неположителен.

Если *средний квадрат* расстояния между парами точек меньше, чем

$$\frac{2nMP}{M-1},$$

то должна существовать пара точек с расстоянием, удовлетворяющим этому неравенству. Каждая точка из этой пары используется с частотой $1/M$. Для разделения этой пары (если бы не было других точек) лучше всего будет служить гиперплоскость, перпендикулярная к середине отрезка, соединяющего эти точки. Каждая точка будет давать вклад в вероятность ошибки, равный вероятности того, что шум смещает точку на половину длины отрезка или еще больше. Таким образом, получаем, что вклад в вероятность ошибки составляет по меньшей мере

$$\begin{aligned} \frac{1}{M} Pr \left\{ \begin{array}{l} \text{смещение от шума в опре-} \\ \text{деленном направлении} \end{array} \right\} &\geq \frac{1}{2} \sqrt{\frac{2nMP}{M-1}} = \\ &= \frac{1}{M} \Phi \left[-\sqrt{\frac{nMP}{(M-1)2N}} \right]. \end{aligned}$$

Этот вклад можно приписать первой из двух рассматриваемых точек. Найденная вероятность ошибки относится к случаю, когда послано

сообщение, соответствующее первой точке, а принятое сообщение находится ближе ко второй и поэтому декодируется как второе или как какое-либо другое сообщение.

Теперь выбросим первое сообщение из этой системы кодовых точек и рассмотрим оставшиеся $(M - 1)$ точек. С помощью тех же доводов можно показать, что среди них существует пара точек, расстояние между которыми больше или равно

$$\sqrt{\frac{2nP(M-1)}{M-2}}.$$

Такая пара дает вклад в вероятность ошибки, связанный с тем, что одна из точек, смещаясь, оказывается ближе к другой, чем к своему первоначальному положению. Этот вклад равен

$$\frac{1}{M} \Phi \left[-\sqrt{\frac{(M-1)nP}{(M-2)2N}} \right].$$

Это рассуждение можно продолжать, последовательно вычеркивая точки и суммируя вклады в вероятность ошибки, пока не останутся только две точки. Таким образом, получим нижнюю границу в следующем виде:

$$P_{e \text{ opt}} \geq \frac{1}{M} \left[\Phi \left(-\sqrt{\frac{nP}{2N} \frac{M}{(M-1)}} \right) + \right. \\ \left. + \Phi \left(-\sqrt{\frac{nP(M-1)}{2N(M-2)}} \right) + \dots + \Phi \left(-\sqrt{\frac{nP}{2N} \frac{2}{1}} \right) \right]. \quad (81)$$

Эту формулу можно несколько упростить, если учесть в ней лишь первые $M/2$ членов (или $(M+1)/2$ членов, если M — нечетное). Так как члены убывают, сумма их только уменьшится, если все члены положить равными последнему. Таким образом, получаем нижнюю границу в виде

$$P_{e \text{ opt}} \geq \frac{1}{2} \Phi \left(-\sqrt{\frac{M}{(M-2)} \frac{nP}{2N}} \right). \quad (82)$$

Для некоторой скорости передачи $R > 0$, когда n возрастает, член $M/(M-2)$ стремится к 1 и граница ведет себя как

$$\frac{1}{2} \Phi \left(-\sqrt{\frac{nP}{2N}} \right).$$

Это асимптотически приближается к

$$\frac{1}{2\sqrt{\frac{\pi nP}{N}}} e^{-(nP)/(4N)}.$$

Отсюда следует, что надежность $E \leq P/(4N) = A^2/4$. Это то же самое значение, что и нижняя граница для E , когда $R \rightarrow 0$.

13. Границы ошибок и другие условия, накладываемые на точки

До сих пор (кроме разд. 12) предполагалось, что все кодовые точки расположены на поверхности сферы, т. е. средний квадрат расстояния между ними равен $\sqrt{nP}$. Рассмотрим задачу оценки вероятности $P'_{e \text{ opt}}(M, n, \sqrt{P/N})$, когда кодовые точки находятся лишь внутри или на самой поверхности сферы. Ясно, что так как при этом ослабляются условия, накладываемые на код, то это приводит только к улучшению, т. е. к уменьшению вероятности ошибки для лучшего кода. Таким образом, $P'_{e \text{ opt}} \leq P_{e \text{ opt}}$.

С другой стороны, покажем, что

$$P'_{e \text{ opt}}\left(M, n, \sqrt{\frac{P}{N}}\right) \geq P_{e \text{ opt}}\left(M, n+1, \sqrt{\frac{P}{N}}\right). \quad (83)$$

В самом деле, предположим, что дан код длиной n , все точки которого расположены на или внутри n -мерной сферы. Добавим к каждому кодовому слову дополнительную координату такой величины, чтобы в $(n+1)$ -мерном пространстве эта точка лежала *в точности* на поверхности $(n+1)$ -мерной сферы. Если первые n координат точки имеют значения $x_1, x_2, \dots, x_n$, удовлетворяющие условию

$$\sum_{i=1}^n x_i^2 \leq nP,$$

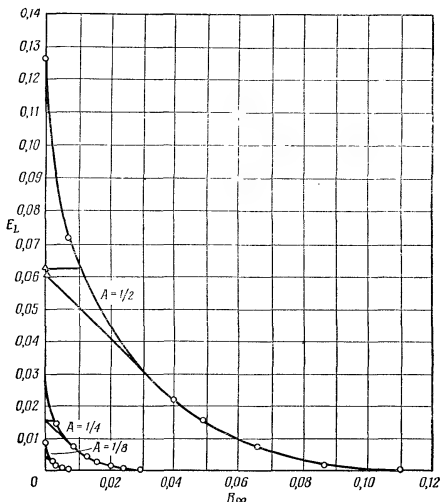
то добавляемая координата должна иметь значение

$$x_{n+1} = \sqrt{(n+1)P - \sum_{i=1}^n x_i^2}.$$

Это приводит к коду первого типа (все точки лежат на поверхности $(n+1)$ -мерной сферы) с M словами длины $n+1$ и отношением сигнал/шум равным P/N . Вероятность ошибки для рассматриваемого кода не больше соответствующей вероятности приведенного кода, поскольку добавление координаты может лишь уточнить процесс декодирования. Можно, например, декодировать, игнорируя последнюю координату, и получать ту же самую вероятность ошибки. Использование этого метода может лишь облегчить декодирование.

Вероятность ошибки для приведенного кода длины $n+1$ должна быть больше или равна аналогичной вероятности для оптимального кода длины $n+1$ с точками на поверхности сферы. Следовательно, имеем выражение (83). Так как $P_{e \text{ opt}}(M, n, \sqrt{P/N})$ при больших n меняется экспоненциально с изменением n , то эффект от замены n на $n+1$ соответствует умножению на постоянную. Поэтому верхние границы $P_{e \text{ opt}}$ остаются без изменения, а ниж-

ние границы оказываются умноженными на величину, которая не сильно зависит от n при большом n . Асимптотические кривые надежности соответственно остаются теми же самыми, благодаря чему приведенные графики кривых E могут применяться в обоих случаях.



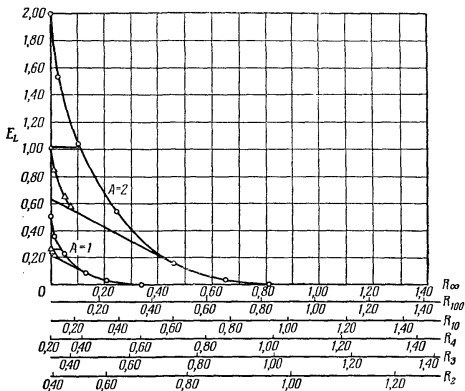
Р и с. 7. Кривые зависимости E_L от R при $A = 1/8, 1/4$ и $1/2$.

Рассмотрим теперь третий тип условий, накладываемых на точки, а именно условие, что *средний* по системе точек квадрат расстояния от начала координат был бы менее или равен nP . Это является опять-таки ослаблением предыдущих условий, и, следовательно, оптимальная вероятность ошибки P''_{opt} будет меньше или равна соответствующей вероятности в предыдущих случаях

$$P''_{\text{opt}}\left(M, n, \frac{P}{N}\right) \leq P'_{\text{opt}}\left(M, n, \frac{P}{N}\right) \leq P_{\text{opt}}\left(M, n, \frac{P}{N}\right). \quad (84)$$

Наши верхние границы вероятности ошибки (и следовательно, нижние границы надежности) можно применять в прежней форме.

Нижние границы $P''_{\epsilon \text{ opt}}$ можно найти следующим образом. Если имеется M точек со средним квадратом расстояния от начала, не превосходящим $P''_{\epsilon \text{ opt}}$, то для некоторого значения α ($0 < \alpha \leq 1$), по крайней мере αM из этих точек лежат внутри поверхности сферы

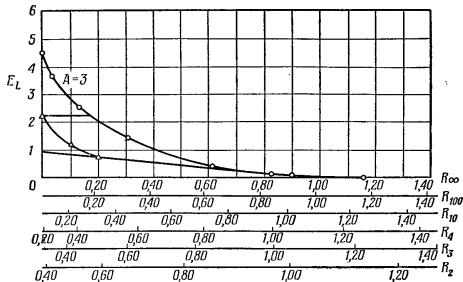


Р и с. 8. Кривые зависимости E_L от R при $A = 1$ и 2 .

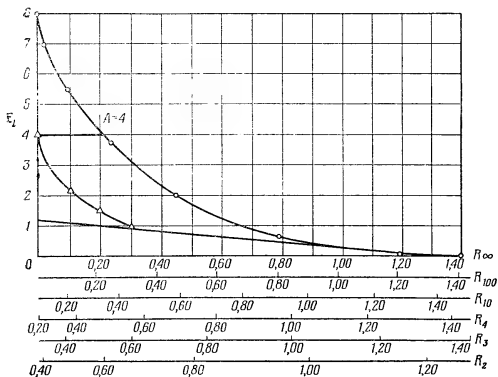
с квадратом радиуса $nP/(1 - \alpha)$. [Если бы более чем $(1 - \alpha)M$ из них находились вне сферы, то они одни внесли бы в сумму квадратов расстояний вклад, больший чем

$$\frac{(1 - \alpha)M_n P}{1 - \alpha},$$

и среднее было бы больше чем nP .] Возьмем оптимальный код, удовлетворяющий третьему условию; так как αM точек его находятся внутри сферы радиуса $\sqrt{nP/(1 - \alpha)}$, из него можно сконструировать код, удовлетворяющий второму условию, с меньшим числом точек и увеличенным радиусом. Вероятность ошибки для нового кода не может превзойти вероятность ошибки первоначального кода больше чем в $1/\alpha$ раз.



Р и с. 9. Кривые зависимости E_L от R при $A = 3$.



Р и с. 10. Кривые зависимости E_L от R при $A = 4$.

(Каждое новое кодовое слово используется самое большее $1/\alpha$ раз. Вероятность ошибки используемого слова каждый раз будет не больше первоначальной.) Поэтому

$$P_{\text{еopt}}^* \left(M, n, \sqrt{\frac{P}{N}} \right) \geq \frac{1}{\alpha} P_{\text{еopt}}^* \left(\alpha M, n, \sqrt{\frac{P}{(1-\alpha)N}} \right) \geq \frac{1}{\alpha} P_{\text{еopt}}^* \left(\alpha M, n+1, \sqrt{\frac{P}{(1-\alpha)N}} \right).$$

14. Кривые для асимптотических границ

Эти кривые вычислены для облегчения нахождения экспонент в асимптотических границах. Основные кривые построены для значений $A = 1/8, 1/4, 1/2, 1, 2, 3, 4, 8, 16$. На рис. 7—11 величины n и E_L

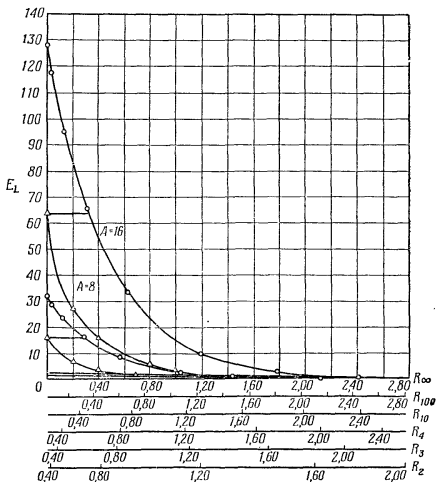
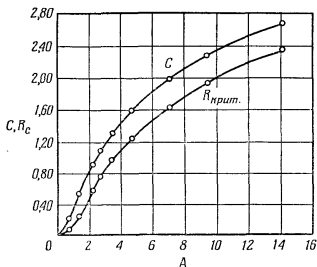
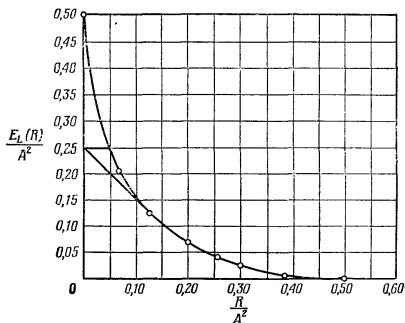


Рис. 11. Кривые зависимости E_L от R при $A = 8$ и 16 .



Р и с. 12. Пропускная способность C и критическая скорость R_c как функции от θ .



Р и с. 13. График $E_L(R)/A^2$ как функции от R/A^2 .

даны в зависимости от скорости передачи R . Так как E_L является функцией θ , а связь между θ и R несколько зависит от n , то потребовалось поместить под кривыми ряд шкал, немного различающихся по R . Это оказалось более удобным способом представления данных, чем использование вспомогательного семейства, связывающего R и θ . Эти же самые кривые дают коэффициент при n в верхней границе (прямая линия вместе с продолжающей ее кривой направо). Точка касания дает критическое R (или критическое θ). Другими словами, кривая и кривая плюс прямая, рассматриваемые в масштабе $n = \infty$, дают верхнюю и нижнюю границы для меры надежности. Верхняя и нижняя границы для E при низких скоростях R также включены в эти кривые. Верхней границей является горизонтальный отрезок, выходящий из точки $R = 0$, $E = a^2/4$. Нижней границей является кривая линия, идущая вниз от этой точки к касательной. Таким образом, значения надежности E лежат в фигуре, ограниченной этими линиями, слева от R_c . Надежность совпадает с кривой справа от R_c . На рис. 12 указана пропускная способность канала C и критическая скорость R_c как функции от θ . Для очень малых A кривая $E_L(R)$ стремится к некоторой предельной кривой. Действительно, если $\theta = (\pi/2) - \varepsilon$, где ε мало, то с помощью очевидных разложений находим, что хорошим приближением является

$$E_L(R) = \frac{A^2}{2} - A\varepsilon + \frac{\varepsilon^2}{2} \quad \text{и} \quad R = \frac{\varepsilon^2}{2}.$$

Исключая ε , получаем

$$\frac{E_L(R)}{A^2} = \frac{1}{2} - \sqrt{\frac{2R}{A^2}}.$$

На рис. 13 изображено отношение $E_L(R)/A^2$ как функция от R/A^2 .

ТЕОРЕМЫ КОДИРОВАНИЯ ДЛЯ ДИСКРЕТНОГО ИСТОЧНИКА ПРИ ЗАДАННОМ КРИТЕРИИ ТОЧНОСТИ¹⁾

Краткое содержание

В статье рассматривается дискретный источник сообщений, образованных последовательностями букв из конечного алфавита. Мера искажения отдельной буквы задается неотрицательной матрицей (d_{ij}). Каждое d_{ij} определяет «цену» или «искажение», возникающее, если буква i воспроизводится в приемнике как буква j . За среднее искажение системой связи (источник — кодирующее устройство — канал с шумами — декодирующее устройство) принимается $d = \sum_{i,j} P_{ij} d_{ij}$, где P_{ij} — вероятность того, что i будет воспроизведено как j . Показано, что существует функция $R(d)$, которая измеряет «эквивалентную скорость» источника для заданного уровня искажения. В том случае когда при кодировании сообщения допускается уровень искажения, равный d , рассматриваемый источник действует подобно источнику со скоростью создания информации $R(d)$. Приведены методы вычисления и рассмотрены различные свойства функции $R(d)$. В заключение развитые теоретические положения обобщаются на эргодические источники, на непрерывные источники и на меры искажения, зависящие от целых блоков букв²⁾.

В статье исследуется проблема кодирования для дискретного источника информации при заданном критерии точности или мере искажения сообщения, восстановленного на приемном конце,

¹⁾ Шаппо К., Coding theorems for discrete source with a fidelity criterion, 1959, *IRE National convention record* (1959), March, 142.

Данное исследование частично субсидировалось министерством ВВС США (научно-исследовательским управлением и авиационным научно-исследовательским командованием, а также министерством ВМФ (морским исследовательским управлением).

²⁾ Аналогичные вопросы, но в более общей ситуации, исследовались также в работах: Колмогоров А. Н., Теория передачи информации, в сб. Серия АН СССР по научн. проблемам автоматизации производства, 1956. Пленарн. заседания, М., АН СССР, 1957, Дискуссия, 148; Пинскер М. С. Вычисление скорости создания сообщений стационарным случайным процессом и пропускной способности стационарного канала, Докл АН СССР (1956) III, 4, 753; Добрушин Р. Л., Общая формулировка основной теоремы Шеннона в теории информации, У.М.Н. (1959) 14, 6, 3.— *Прим. ред.*

относительно фактически переданного сообщения. В ряде случаев может существовать некоторый допустимый уровень искажения, определяемый этой мерой, и нам желательно так кодировать информацию, чтобы достигнуть максимально возможной скорости передачи, не превосходя этого допустимого уровня искажения. Данная работа представляет собой обобщение и детализацию высказанных ранее идей с подробным рассмотрением случая дискретного канала ¹⁾.

Покажем, что для широкого класса мер искажения и дискретных источников информации существует функция $R(d)$ (зависящая от заданных меры искажения и источника), которая измеряет в определенном смысле эквивалентную скорость R источника (битах на букву), когда d — допустимый уровень искажения. Будут даны методы вычисления точных значений $R(d)$ в некоторых простых случаях и асимптотических значений $R(d)$ в более сложных случаях. Основные результаты работы, грубо говоря, показывают, что по каналу без памяти с пропускной способностью C (бит в секунду) и с уровнем искажения меньшим или равным d , невозможно передавать сигналы со скоростью, большей чем $C/R(d)$ (букв в секунду), но при достаточно длинных блоковых кодах возможно сколь угодно точно приблизиться к скорости $C/R(d)$ с уровнем искажения d .

Наконец, приведены детально разобранные частные примеры, в которых используются вероятность искажения буквы сообщения и некоторые другие простые меры уровня искажения.

Мера искажения отдельной буквы

Предположим, что имеется дискретный источник информации, вырабатывающий последовательность букв или «слово» $m = m_1, m_2, \dots, m_i, \dots$; каждая буква выбрана из конечного алфавита. Они должны быть переданы и воспроизведены, хотя бы приближенно, на приемном конце. Пусть $Z = z_1, z_2, \dots, z_i, \dots$ — воспроизведенное слово. Буквы z_i могут быть как из того же самого алфавита, что и буквы m_i , так и из расширенного алфавита, содержащего специальные символы для нераспознанных или не вполне распознанных букв.

В случае телеграфирования при наличии помех m и Z могут, например, иметь следующий вид:

$m = \text{I HAVE HEARD THE MERMAIDS SINGING...}^2)$

$Z = \text{I H?VT HEA?D TSE B?RMAIDZ ??NGING...}$

При этом алфавит Z состоит из обычных букв и пробела, образующих алфавит m , и дополнительных символов «?», « $\hat{A}$ », « $\hat{B}$ » и т. д.,

¹⁾ В е п н е т W., New results in the calculation of modulation products, *BSTJ* (1933), April, 228.

²⁾ Я слышал пение русалок. — Прим. ред.

указывающих на некоторую неопределенность в отождествлении. Иногда алфавит Z может целиком отличаться от алфавита m .

Рассмотрим случай, когда задана мера точности передачи или «искажения», измеряющая расхождение между переданными и принятыми словами. Сначала исследуем специальный вид меры искажения, который явится основой для существенных обобщений.

Определим меру искажения отдельной буквы. Пусть задана матрица (d_{ij}) с $d_{ij} \geq 0$. Здесь i принимает значения, соответствующие алфавиту m , состоящему, скажем, из a букв (предполагается, что буквы перенумерованы), в то время как j принимает значения, соответствующие алфавиту Z . Величину d_{ij} можно рассматривать как стоимость воспроизведения буквы i буквой j . Пусть алфавит Z включает в себя алфавит m . Тогда будем считать, что искажение между буквой алфавита m и ее точным воспроизведением равно нулю и при любом неточном воспроизведении положительно. В этом случае удобно использовать одинаковую индексацию обоих алфавитов, так что $d_{ii} = 0$, $d_{ij} > 0$ ($i \neq j$).

Если слово m воспроизведено как слово Z , то искажение d измеряется следующим образом:

$$d(m, Z) = \frac{1}{t} \sum_{k=1}^t d_{m_k} z_k.$$

Если в системе связи слово m встречается с вероятностью $P(m)$ и если условная вероятность того, что при передаче слово m будет воспроизведено как слово Z , равна $P(Z|m)$, то полное искажение системы определяется выражением

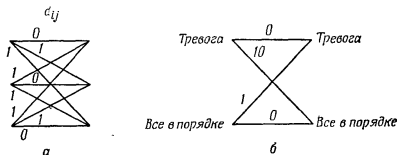
$$d = \sum_{m, Z} P(m) P(Z|m) d(m, Z).$$

Здесь предполагается, что все сообщения и воспроизведенные слова имеют одну и ту же длину t . В системах кодирования с переменной длиной слов аналогичной мерой является просто сумма по i и j произведений d_{ij} на полную вероятность того, что буква i воспроизведена как буква j . Заметим, что $d = 0$ тогда и только тогда, когда каждое слово с вероятностью единица воспроизводится точно; в противном случае $d > 0$ (в случаях, когда алфавит Z содержит алфавит m).

Несколько простых примеров

Меру искажения можно задавать при помощи матрицы (d_{ij}) , все элементы которой неотрицательны. Другим способом представления меры искажения может служить линейная схема, подобная схемам, используемым для изображения канала без памяти с шумом. Но при этом линиям схемы вместо вероятностей должны сопоставляться величины d_{ij} .

Вероятность ошибки на букву представляет собой простой пример меры искажения при совпадающих алфавитах m и Z . В этом случае при одинаковом упорядочении алфавитов $d_{ij} = 1 - \delta_{ij}$ ¹⁾. На рис. 1, а приведен соответствующий график для случая, когда алфавиты m и Z содержат по 3 буквы. Такая мера искажения может быть пригодна для оценки точности работы телетайпа или системы дистанционного набора.



Р и с. 1.

Другим примером является передача квантованных положений штурвала или ручки управления. Предположим, что окружность разделена на пять равных дуг. Ошибку в ту или другую сторону

	$-\frac{1}{2} + \frac{1}{2}$	
-1	1	2
0	1	1
1	2	1

Р и с. 2.

на один сегмент будем расценивать $1/2$, а большие ошибки — единицей. Тогда мера искажения может быть представлена в виде

$$d_{ij} = \begin{cases} 0, & i = j, \\ 1/2, & |i - j| = 1, \text{ mod } 5, \\ 1, & |i - j| > 1, \text{ mod } 5. \end{cases}$$

Третьим примером может служить двоичная система, ежесекундно передающая информацию об одном из двух состояний, «все в порядке» или «тревога». Вообще говоря, значительно важнее, чтобы правильно был принят сигнал «тревога», чем сигнал «все в порядке».

¹⁾ Здесь $\delta_{ij} = \begin{cases} 1 & \text{при } i = j, \\ 0 & \text{при } i \neq j. \end{cases}$ — Прим. ред.

Если соответствующим сигналам приписать веса 10 и 1, то схема будет иметь вид, показанный на рис. 1, б.

Четвертым примером, в котором алфавиты m и Z целиком различны, является случай, когда алфавит m содержит три возможных отсчета показаний $-1, 0, +1$, а воспроизводимый алфавит, возможно, по каким-либо соображениям экономии состоит из двух букв $-1/2$ и $+1/2$. Соответствующая матрица показана на рис. 2.

Скорость при заданном искажении $R(d)$

Предположим теперь, что последовательные буквы сообщения статистически независимы, но выбираются всякий раз с одними и теми же вероятностями. P_i — вероятность буквы i из алфавита. Такой тип источника назовем *источником с независимыми буквами*.

Имея такую совокупность вероятностей P_i и меру искажения d_{ij} , можно следующим образом определить *скорость при заданном искажении*. Выберем произвольное множество вероятностей $q_i(j)$ перехода от i к j . (Конечно, $q_i(j) \geq 0$ и $\sum_j q_i(j) = 1$). При этом можно

вычислить две величины: меру искажения $d(q_i(j)) = \sum_{i,j} P_i q_i(j) d_{ij}$, когда буква i воспроизводится как j с условной вероятностью $q_i(j)$, и среднюю взаимную информацию между i и j для того же случая, а именно

$$R(q_i(j)) = E \log \frac{q_i(j)}{\sum_k P_k q_k(j)} = \sum_{i,j} P_i q_i(j) \log \frac{q_i(j)}{\sum_k P_k q_k(j)}.$$

Скорость $R(d^)$ при заданном искажении d^* определяется как нижняя грань $R(q_i(j))$ при варьировании $q_i(j)$ таким, чтобы соблюдались естественные вероятностные ограничения и среднее искажение d оставалось меньшим или равным d^* .*

Заметим, что $R(q_i(j))$ — непрерывная функция от $q_i(j)$ в области изменения $q_i(j)$. Эта область замкнута. Следовательно, нижняя грань R действительно достигается и является минимумом всевозможных значений R . Далее, из этого определения ясно, что $R(d)$ — монотонно убывающая функция d .

Выпуклость вниз кривой $R(d)$

Предположим, что на кривой $R(d)$ имеются две точки: (R, d) — достигнутая при выборе в качестве вероятностей перехода $q_i(j)$ и (R', d') — достигнутая при выборе $q'_i(j)$. Рассмотрим линейную комбинацию этих вероятностей $q''_i(j) = \lambda q_i(j) + (1 - \lambda) q'_i(j)$. Такая линейная комбинация приводит к значению d''_i не большему (в силу линейности d) чем $\lambda d + (1 - \lambda) d'$. С другой стороны, известно,

что $R(q_i(j))$ — скорость передачи по каналу — является выпуклой функцией переходных вероятностей канала¹⁾. Отсюда $R'' \leq \lambda R + (1 - \lambda) R'$. Минимизация R по $q_i(j)$ при фиксированном d'' должна привести к величине, не превосходящей R'' . Поэтому кривая R как функция от d (или наоборот) выпукла вверх.

Очевидно, что d принимает минимально возможное значение, если при каждом i переходной вероятности $q_i(j)$ приписано значение 1 для тех j , при которых d_{ij} минимально. Таким образом, наименьшее возможное d можно записать как

$$d_{\min} = \sum_i P_i \min_j d_{ij}.$$

Так если алфавит m отображается в алфавит Z , то $d_{\min} = 0$ и соответствующее значение R есть обычная энтропия или скорость создания сообщения источником. В более общем случае, если только существует единственный $\min_j d_{ij}$, то $R(d_{\min})$ может быть также легко

найден. Для этого R надо вычислить при упомянутом выше задании $q_i(j)$. В других случаях вычисление $R(d_{\min})$ несколько сложнее.

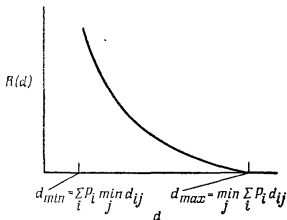
С другой стороны, так как средняя взаимная информация положительна, за исключением случая, когда события независимы, то $R = 0$ тогда и только тогда, когда $q_i(j) = Q_j$ является функцией только j . Для заданных Q_j , определяющих $R = 0$, искажение d равно тогда $\sum_{i,j} P_i Q_j d_{ij} = \sum_j Q_j \sum_i P_i d_{ij}$. Сумма по i в последнем выражении неотрицательна. Для того чтобы минимизировать d при $R = 0$, надо найти j (скажем j^*), при котором $\sum_i P_i d_{ij}$ имеет минимум, и принять, что $Q_{j^*} = 1$. Это можно осуществить, положив $q_i(j^*) = 1$ (все другие значения $q_i(j)$ равны 0).

Итак, $R(d)$ есть выпуклая вниз функция, изменяющаяся от $R(d_{\min})$ в точке $d_{\min} = \sum_i P_i \min_j d_{ij}$ до 0 в точке $d_{\max} = \min_j \sum_i P_i d_{ij}$, как показано на рис. 3. Внутри этого интервала в силу выпуклости²⁾ имеет место непрерывность в обе стороны (R как функция от d и d как функция от R) и функция строго монотонно убывает. Для $d \geq d_{\max}$ скорость $R = 0$. Легко также видеть, что для нахождения $R(d^*)$ при произвольном d^* из интервала от $d_{\min}$ до $d_{\max}$ надо задавать $q_i(j)$ так, чтобы d удовлетворяло равенству $d = d^*$ (а не неравенству $d < d^*$). При $d^* > d_{\max}$ для $q_i(j)$, минимизирующего R , может иметь место и неравенство. Следовательно, проблема минимизации может быть ограничена рассмотрением минимума в подпространстве $d = d^*$, исключая область $d^* > d_{\max}$ (где $R(d^*) = 0$).

¹⁾ Это нетрудно проверить прямым вычислением. — Прим. ред.

²⁾ И монотонности. — Прим. ред.

Выпуклость R как функции заданных $q_i(j)$ помогает в вычислении $R(d)$ в специальных случаях. Из этого вытекает, что любой локальный минимум (в подпространстве, соответствующем фиксированному d) есть абсолютный минимум в этом подпространстве. В противном случае можно было бы соединить локальный и абсолютный минимумы отрезком прямой и найти вдоль этого отрезка непрерывный ряд точек, лежащих ниже локального минимума. Это противоречит определению локального минимума.



Р и с. 3.

Далее, функции $R(q_i(j))$ и $d(q_i(j))$ имеют непрерывные производные внутри допустимого множества значений $q_i(j)$. Поэтому для нахождения минимума могут быть использованы обычные вычислительные методы (например, множители Лагранжа). Но, вообще говоря, при этом возникает необходимость решения системы совместных уравнений.

Нахождение $R(d)$ в некоторых простых случаях

Один частный случай приводит к простому и явному выражению для $R(d)$. Предположим, что все a букв исходного сообщения равновероятны: $P_i = 1/a$. Далее, предположим, что матрица (d_{ij}) — квадратная и такая, что все строки состоят из одного и того же множества элементов и все столбцы также содержат одно и то же множество элементов, хотя, конечно, в различном порядке.

Примером такого случая может служить упомянутый выше штурвал, если все его положения равновероятны. Другим примером является случай, когда все буквы равновероятны и мера искажения есть просто вероятность ошибки.

В общем случае пусть элементы любой строки или столбца будут $d_1, d_2, \dots, d_a$. Тогда можно показать, что для заданного d минимум R

будет иметь место, когда всем линиям схемы, которым соответствуют искажения d_k , приписаны вероятности

$$q_k = \frac{e^{-\lambda d_k}}{\sum_i e^{-\lambda d_i}}.$$

Здесь λ — параметр, изменяющийся от 0 до ∞ и определяющий значение величины d . Для этого минимизирующего значения величины d и R выражаются через параметр λ :

$$d = \frac{\sum_i d_i e^{-\lambda d_i}}{\sum_i e^{-\lambda d_i}},$$

$$R = \log \frac{a}{\sum_i e^{-\lambda d_i}} - \lambda d.$$

Легко видеть, что когда $\lambda = 0$, то $d = \frac{1}{a} \sum_i d_i$ и $R = 0$. Если $\lambda \rightarrow \infty$, то $d \rightarrow d_{\min}$ и $R \rightarrow \log \frac{a}{k}$, где k — число d_i , равных $d_{\min}$.

Это может быть доказано следующим образом. Предположим, что нам заданы $q_i(j)$, определяющие некоторые d^* и R^* . Теперь для всех линий схемы с $d_{ij} = d_i$ зададим новое значение вероятности перехода $q_i(j)$, определяемое как среднее из заданных первоначально для этих линий значений $q_i(j)$. Аналогично для каждой линии с $d_{ij} = d_2$ $q_i(j)$ задается как среднее из первоначально заданных для этих линий значений $q_i(j)$ и т. д. В силу линейности d величина d при этих новых $q_i(j)$ не изменяет своего значения, равного d^* . Покажем, что новое R останется равным или будет меньше чем R^* . Величину R можно записать как $H(m) - H(m|Z)$. Для новых $q_i(j)$ $H(m)$ не изменяется, а $H(m|Z)$ может только возрасти, так как в силу выпуклости функции $-\sum_i x_i \log x_i$ имеем

$$-\sum_j \alpha_j \sum_i x_j^{(t)} \log x_j^{(t)} \geq -\sum_i \left(\sum_j \alpha_j x_j^{(t)} \right) \log \sum_j \alpha_j x_j^{(t)},$$

где при заданном t $x_j^{(t)}$ — совокупность вероятностей и α_j — совокупность весовых множителей. В частности,

$$-\sum_j \frac{\sum_s q_j^{(s)}}{\sum_{s,i} q_i^{(s)}} \sum_i \frac{q_j^{(t)}}{\sum_s q_j^{(s)}} \log \frac{q_j^{(t)}}{\sum_s q_j^{(s)}} \geq -\sum_i \frac{\sum_j q_j^{(t)}}{\sum_{s,i} q_i^{(s)}} \log \frac{\sum_j q_j^{(t)}}{\sum_{s,i} q_i^{(s)}},$$

где $q_j^{(s)}$ есть первоначальное значение вероятности перехода соответствующей линии, такой, что она идет от буквы s и ей сопо-

ставлено d_j . Но левая часть этого неравенства может быть интерпретирована как $H(m|Z)$ после операции усреднения, в то время как правая часть есть $H(m|Z)$ до усреднения. Отсюда следует искомый результат.

Итак, при минимизирующих значениях $q_i(j)$ все линии с одинаковыми d будут иметь одинаковые переходные вероятности. Будем теперь обозначать через q_i переходную вероятность линий, которым сопоставлено d_i . Скорость R и искажение d могут быть теперь записаны в виде

$$d = \sum_i q_i d_i,$$

$$R = \log a + \sum_i q_i \log q_i,$$

так как все буквы алфавита Z теперь равновероятны и

$$H(m) = \log a, \quad H(m|Z) = - \sum_i q_i \log q_i.$$

Далее нам надо так выбрать q_i , удовлетворяющие условию $\sum_i q_i = 1$, чтобы при заданном d минимизировать R . Для этого применим множители Лагранжа:

$$U = \log a + \sum_i q_i \log q_i + \lambda \sum_i q_i d_i + \mu \sum_i q_i,$$

$$\frac{dU}{dq_i} = 1 + \log q_i + \lambda d_i + \mu = 0,$$

$$q_i = A e^{-\lambda d_i}.$$

Выбирая $A = \frac{1}{\sum_i e^{-\lambda d_i}}$, удовлетворим условию $\sum_i q_i = 1$ и полу-

чим таким образом стационарную точку. В силу выпуклости кривой $R(d)$ эта стационарная точка должна быть абсолютным минимумом R для соответствующего значения d . Подставляя это значение вероятности в формулы для d и R , получим приведенные выше выражения.

Скорость передачи для произведения источников с мерой, являющейся суммой мер искажения

Предположим, что имеется два независимых источника, каждый со своей собственной мерой искажения d_{ij} и d'_{ij} , и скоростями передачи $R_1(d_1)$ и $R_2(d_2)$. Предположим также, что каждый источник создает одну букву в секунду. Рассмотрим упорядоченные пары букв как отдельные буквы, создаваемые комбинированным источ-

ником, который будем называть *произведением источников*. Если общее искажение измеряется суммой отдельных искажений $d = d_1 + d_2$, то существует простой метод определения функции $R(d)$ для произведения источников. В самом деле, покажем, что $R(d)$ получается суммированием координат кривых $R_1(d_1)$ и $R_2(d_2)$ в точках, где обе кривые имеют одинаковый наклон, и что переходные вероятности, соответствующие произвольной точке $R(d)$, являются произведением переходных вероятностей для точек-компонент.

Покажем сначала, что минимум R для произведения источников, взятый по всевозможным допустимым значениям $q_{i'j'}(j, j')$, не увеличится, если операцию минимизации провести по переходным вероятностям, заданным в виде произведения $q_i(j) q_{i'}(j')$, где q и q' выражаются через $q_{i'j'}(j, j')$ следующим образом:

$$q_i(j) = \sum_{i', j'} P_{i'} q_{i'j'}(j, j'),$$

$$q_{i'}(j') = \sum_{i, j} P_i q_{i'j'}(j, j').$$

Поскольку эти выражения неотрицательны и суммирование по j и j' соответственно дает единицу, то эти выражения удовлетворяют ограничениям, накладываемым на переходные вероятности. Кроме того, переходные вероятности $q_i(j) q_{i'}(j')$ дают точно такое общее искажение, что и переходные вероятности $q_{i'j'}(j, j')$. Действительно,

$$\begin{aligned} \sum_{i, i', j, j'} P_i P_{i'} q_i(j) q_{i'}(j') (d_{ij} + d_{i'j'}) &= \sum_{i, j} P_i q_i(j) d_{ij} + \sum_{i', j'} P_{i'} q_{i'}(j') d_{i'j'} = \\ &= \sum_{i, i', j, j'} P_i P_{i'} q_{i'j'}(j, j') (d_{ij} + d_{i'j'}), \end{aligned}$$

где последний член можно рассматривать как искажение при переходных вероятностях $q_{i'j'}(j, j')$.

Теперь, если вместо $q_i(j) q_{i'}(j')$ взять $q_{i'j'}(j, j')$, то взаимная информация R может или уменьшиться или остаться постоянной. В самом деле, эта средняя взаимная информация может быть записана в терминах энтропии следующим образом: отмечаем звездочками значения энтропии, соответствующие переходным вероятностям $q_i(j) q_{i'}(j')$ и опускаем звездочку для значений энтропии, соответствующих $q_{i'j'}(j, j')$. Имеем

$$\begin{aligned} R = H(i, i') - H(i, i' | j, j') &\geq H(i, i') - H(i' | j) - H(i' | j') = \\ &= H(i, i') - H^*(i | j) - H^*(i' | j'). \end{aligned}$$

(Здесь используется тот факт, что при нашем определении $q_i(j)$ и $q_{i'}(j')$ справедливы равенства $Pr^*(i | j) = Pr(i | j)$ и $Pr^*(i' | j') = Pr(i' | j')$, которые становятся очевидными, если выписать соответствующие вероятности.) Последний член предыдущей формулы

равен R^* , так как в силу независимости источников $H(i, i') = H(i) + H(i') = H^*(i) + H^*(i')$.

Таким образом, наше предварительное утверждение доказано. Из него следует, что любая точка на кривой $R(d)$ для произведения источников получается с помощью произведения переходных вероятностей (как при независимых событиях) и, следовательно, есть сумма координат соответствующих точек обеих кривых.

Очевидно, что наилучший выбор R при заданном искажении d дается выражением

$$R(d) = \min_t [R_1(t) + R_2(d-t)],$$

и этот минимум будет иметь место, когда

$$\left| \frac{d}{dt} R_1(t) \right| = \left| \frac{d}{dt} R_2(d-t) \right|.$$

Итак, точки компонент, которые должны суммироваться, являются точками, где составляющие кривые имеют одинаковую кривизну. Выпуклость этих кривых гарантирует единственность такой пары точек для любого d .

Нижняя граница для искажения при заданной пропускной способности канала

Значение функции $R(d)$ состоит в том, что она определяет пропускную способность канала, требуемую для того, чтобы передавать сообщения с определенной скоростью и с некоторым минимальным искажением. Рассмотрим следующий случай. Пусть задан источник независимых букв с вероятностями P_i появления различных возможных букв и пусть задана мера искажения d_{ij} отдельной буквы, причем d соответствует скорости $R(d)$. Наконец, имеется дискретный канал K без памяти с пропускной способностью C бит в секунду (предполагается, что канал используется один раз каждую секунду). Требуется передавать через канал посредством блокового кода слова длины t . Длина кодовых слов в канале равна n . Каково наименьшее искажение d , которое может быть достигнуто с помощью системы кодирования и декодирования такого типа?

Теорема 1. При сделанных выше предположениях не существует кода, обеспечивающего искажение d меньшее, чем (минимальное) d^* , удовлетворяющее условию

$$R(d^*) = \frac{n}{t} C,$$

или, что равносильно, для любого кода $d \geq \varphi\left(\frac{n}{t} C\right)$, где φ — функция, обратная $R(d)$.

Эта теорема и обратный ей позитивный результат, который будет дан ниже, показывают, что $R(d)$ можно рассматривать как эквивалентную скорость источника при заданном искажении d .

Теорема 1 утверждает, что при искажении d и при t буквах текста общее число битов, подаваемых в канал, равное $tR(d)$, не должно превосходить пропускной способности при n -кратном использовании канала. Обратная теорема показывает, что, выбирая достаточно большие n и t и подходящие коды, возможно приблизиться как угодно близко к этой предельной пропускной способности.

Для доказательства¹⁾ теоремы 1 предположим, что задан блочный код, которым кодируются все слова длины t в сигналы длины n , а операция декодирования состоит в отображении сигнала длины n на выходе канала в слова алфавита Z длины t . Пусть слово сообщения представляется в виде $m = m_1, m_2, \dots, m_t$, соответствующее слово на входе канала — в виде $X = x_1, x_2, \dots, x_n$, слово на выходе канала — в виде $Y = y_1, y_2, \dots, y_n$, а воспроизводимое слово — в виде $Z = z_1, z_2, \dots, z_t$. При данных кодирующей и декодирующей системах X есть функция m , а Z — функция Y . Буквы m_i выбираются независимо друг от друга в соответствии с заданными вероятностями букв, а переходные вероятности в канале определяют множество условных вероятностей $P(y|x)$, отнесенных к каждой паре x_i, y_i . Наконец, источник и канал независимы в том смысле, что $P(Y|m, X) = P(Y|X)$.

Прежде всего покажем, что $H(m|Z) \geq H(m) - nC$. Имеем $H(m|Z) \geq H(m|Y)$ (так как Z — есть функция Y), а также $H(m|Y) = H(X|Y) - H(X) + H(m)$. Последнее следует из приведенного выше условия независимости. Действительно, $H(Y|m, X) = H(Y|X)$, так что $H(Y, m, X) - H(m, X) = H(X, Y) - H(X)$. Но $H(m, X) = H(m)$, так как X есть функция m и по той же причине $H(m, X, Y) = H(m, Y)$. Упорядочивая приведенные соотношения, имеем

$$\begin{aligned} H(X, Y) &= H(m, Y) + H(X) - H(m, X) = \\ &= H(m, Y) + H(X) - H(m), \\ H(X|Y) &\leq H(m|Y) + H(X) = H(m). \end{aligned}$$

Здесь мы, воспользовались тем, что $H(m, X) = H(m)$ и затем вычли выражение $H(Y)$ из каждой части равенства. Отсюда получаем $H(m|Z) \geq H(X|Y) - H(X) + H(m)$.

¹⁾ Доказательство этой теоремы может быть проведено проще, см. работы: Добрушин Р. Л., Общая формулировка основной теоремы Шеннона в теории информации, *У. М. Н.* (1959) 14, 6, 3, и Добрушин Р. Л., Передача информации по каналу с обратной связью, *Теория вероятностей и ее применение* (1958), 3, 4, 395. — Прим. ред.

Теперь покажем, что $H(X|Y) \geq nC$. Для этого воспользуемся методом, который уже использовали в других подобных случаях, а именно рассмотрим изменение величины $H(X|Y)$ с каждой принятой буквой. Итак (используя Y_k для обозначения первых k букв сигнала на выходе и т. д.),

$$\begin{aligned} \Delta H(X|Y) &= H(X|y_1, y_2, \dots, y_k) - H(X|y_1, y_2, \dots, y_{k+1}) = \\ &= H(X, Y_k) - H(Y_k) - H(X, Y_k, y_{k+1}) + H(Y_k, y_{k+1}) = \\ &= H(y_{k+1}|Y_k) - H(y_{k+1}|X, Y_k) = H(y_{k+1}|Y_k) - H(y_{k+1}|x_{k+1}) \leq \\ &\leq H(y_{k+1}) - H(y_{k+1}|x_{k+1}) \leq C. \end{aligned}$$

Здесь использован тот факт, что рассматривается канал без памяти, так что $P(y_{k+1}|X, Y_k) = P(y_{k+1}|x_{k+1})$ и, следовательно, $H(y_{k+1}|X, Y_k) = H(y_{k+1}|x_{k+1})$. Наконец, последнее неравенство обусловлено тем, что C — наибольшее из возможных значений $H(y) - H(y|x)$.

Так как элементарное приращение величины $H(X|Y_k)$ ограничено значением C , то полное изменение после n шагов ограничено величиной nC . Следовательно, величина $H(X|Y)$ не меньше $H(X) - nC$. Отсюда

$$\begin{aligned} H(m|Z) &\geq H(X|Y) - H(X) + H(m) \geq \\ &\geq H(X) - nC - H(X) + H(m), \\ H(m|Z) &\geq H(m) - nC. \end{aligned} \quad (1)$$

Теперь найдем верхнюю границу $H(m|Z)$, выраженную в терминах искажения d . Имеем

$$\begin{aligned} H(m|Z) &= H(m_1, m_2, \dots, m_t|z_1, z_2, \dots, z_t) \leq \sum_i H(m_i|z_i) = \\ &= \sum_i H(m_i) - \sum_i [H(m_i) - H(m_i|z_i)]. \end{aligned}$$

Величина $H(m_i) - H(m_i|z_i)$ есть средняя взаимная информация между буквой m_i исходного сообщения и воспроизведенной буквой z_i . Если обозначим через d_i искажение между этими буквами, то скорость $R(d_i)$ (соответствующая этому d_i) удовлетворяет условию

$$R(d_i) \leq H(m_i) - H(m_i|z_i),$$

так как $R(d_i)$ есть минимум взаимной информации при искажении d_i . Поэтому наше неравенство можно записать в виде

$$H(m|Z) \leq \sum_{i=1}^t H(m_i) - \sum_{i=1}^t R(d_i).$$

Теперь, используя то обстоятельство, что $R(d)$ выпуклая вниз функция, имеем

$$H(m|Z) \leq \sum_i H(m_i) - tR\left(\sum_i \frac{d_i}{t}\right).$$

Но $\sum_i d_i/t = d$ есть полное искажение системы, так что

$$H(m|Z) \leq \sum_i H(m_i) - tR(d).$$

Комбинируя это неравенство с неравенством (1) и используя предположение о независимости букв, из которого следует, что $H(m) = \sum_i H(m_i)$, имеем

$$\begin{aligned} H(m) - nC &\leq H(m) - tR(d), \\ nC &\geq tR(d). \end{aligned}$$

Это и есть по существу утверждение, сформулированное в теореме 1.

Следует заметить, что доказанная теорема представляет собой утверждение о минимуме искажения после любого конечного числа n использований канала. Это точный результат, а не асимптотика, справедливая для больших n . Итак, как видно из метода доказательства, для любого блочного или неравномерного кода воспроизведение t или более букв на приемном конце, какова бы ни была принятая последовательность, возможно только после n -кратного использования канала.

Теоремы кодирования при заданной мере искажения отдельной буквы

Докажем теперь позитивную теорему кодирования, соответствующую негативным утверждениям теоремы 1, а именно докажем, что возможно приблизиться к нижней границе искажений при заданном отношении числа n букв сигнала к t буквам сообщения. Рассмотрим эргодический источник сообщений с мерой искажения отдельной буквы d_{ij}^1). Выбор такого более общего источника вместо источника с независимыми буквами, рассмотренного в теореме 1, окажется полезным при последующем обобщении теоремы. И для эргодического источника будут иметься, конечно, вероятности отдельных букв P_i . Тогда можно определить $R(d)$ через эти вероятности точно так же, как в случае источника независимых букв.

Сначала докажем следующую лемму.

¹⁾ То есть источник такой, что случайный процесс, описывающий его, является эргодическим. — Прим. ред.

Лемма 1. Предположим, что имеется эргодический источник с вероятностями букв P_i , мерой искажения отдельных букв d_{ij} и с заданным множеством переходных вероятностей $q_i(j)$, так что

$$\sum_{i,j} P_i q_i(j) d_{ij} = d^*,$$

$$\sum_{i,j} P_i q_i(j) \log \frac{q_i(j)}{\sum_k P_k q_k(j)} = R.$$

Пусть $Q(Z)$ — вероятностная мера последовательности Z из пространства воспроизводимых последовательностей, получающаяся, если последовательные буквы источника имеют вероятности перехода $q_i(j)$ независимо от предыдущих и последующих букв в буквы алфавита Z^1). Тогда при данном $\varepsilon > 0$ для всех достаточно больших длин блоков t существует множество α сообщений длины t с общей вероятностью $P(\alpha) \geq 1 - \varepsilon$ и для каждого m , принадлежащего α , множество Z блоков длины t , скажем β_m , такое, что

$$1) \quad d(m, Z) \leq d^* + \varepsilon \quad \text{для} \quad m \in \alpha \text{ и } Z \in \beta_m,$$

$$2) \quad Q(\beta_m) \geq e^{-t(R+\varepsilon)} \quad \text{для любого} \quad m \in \alpha.$$

Другими словами, говоря не совсем строго, длинные сообщения будут с высокой вероятностью попадать в некоторое подмножество α . Каждому элементу m этого подмножества сопоставлено множество β_m последовательностей Z . Искращения между m и элементами множества β_m могут быть (самое худшее) лишь немногим больше, чем d^* , и логарифм полной вероятности по мере Q множества β_m ограничен величиной $e^{-t(R+\varepsilon)}$.

Для доказательства леммы рассмотрим блоки сообщения длины t и блоки длины t в алфавите Z и рассмотрим две случайные величины: искажение d между m -блоком и Z -блоком и (неусредненную) взаимную информацию, имеющие вид

$$d = \frac{1}{t} \sum_i d_{m_i z_i},$$

$$I(m, Z) = \frac{1}{t} \log \frac{Pr(Z|m)}{Q(Z)} = \frac{1}{t} \sum_i \log \frac{Pr(z_i|m_i)}{Q(z_i)}.$$

Здесь m_i — буква m -блока и z_i — буква Z -блока. Как d , так и R — случайные величины, принимающие различные значения, соответствующие различным выборам m и Z . d и R — суммы t случайных величин, являющихся функциями значений совместного процесса m, Z , идентичны друг другу с точностью до сдвига на t позиций.

Так как совместный процесс эргодичен, то можно применить

¹⁾ При построении меры $Q(Z)$ предполагается также, что она вычисляется так, как если бы буквы источника были независимы. — Прим. ред.

эргодическую теорему и утверждать, что, когда t велико, d и R с вероятностью, почти равной 1, близки к своим средним значениям. В частности, для любых данных ε_1 и δ , если t достаточно велико, то с вероятностью, большей или равной $1 - \delta^2/2$, будем иметь неравенство

$$d \leq \sum_{i,j} P_i q_i(j) d_{ij} + \varepsilon_1 = d^* + \varepsilon_1.$$

Точно так же с вероятностью, не меньшей $1 - \delta^2/2$, будем иметь неравенство

$$I \leq \sum_{i,j} P_i q_i(j) \log \frac{q_i(j)}{Q_j} + \varepsilon_1 = R(d^*) + \varepsilon_1.$$

Пусть γ — множество пар (m, Z) , для которых верны оба неравенства. Тогда $Pr(\gamma) \geq 1 - \delta^2$, так как каждое из условий может исключать самое большее множество вероятности $\delta^2/2$. Теперь для любого m_1 определим β_{m_1} , как множество Z , такое, что (m_1, Z) принадлежит γ ; пусть α есть множество, образованное из m , для которых справедливо неравенство $Pr(\beta_m | m) \geq 1 - \delta$. Вероятность множества α должна удовлетворять условию $Pr(\alpha) \geq 1 - \delta$, так как в противном случае полная вероятность множества, дополнительного к γ , будет не меньше $\delta \cdot \delta = \delta^2$: первое δ было бы вероятностью того, что m не принадлежит α , второе δ — условной вероятностью того, что для таких m последовательность букв Z не принадлежит β_m , а произведение их дает нижнюю границу вероятности множества, дополнительного к γ . Это приводит к противоречию.

Если $Z \in \beta_{m_1}$, то

$$\begin{aligned} \frac{1}{t} \log \frac{Pr(Z | m_1)}{Q(Z)} &\leq R(d^*) + \varepsilon_1, \\ Pr(Z | m_1) &\leq Q(Z) e^{t(R(d^*) + \varepsilon_1)}, \\ Q(Z) &\geq Pr(Z | m_1) e^{-t(R(d^*) + \varepsilon_1)}. \end{aligned}$$

Суммируя эти неравенства по всем $Z \in \beta_{m_1}$, получаем

$$Q(\beta_{m_1}) = \sum_{Z \in \beta_{m_1}} Q(Z) \geq e^{-t(R + \varepsilon_1)} \sum_{Z \in \beta_{m_1}} Pr(Z | m_1).$$

Если $m_1 \in \alpha$, то, как показано выше, $\sum_{Z \in \beta_{m_1}} Pr(Z | m_1) \geq 1 - \delta$ и неравенство может быть приведено к следующему виду:

$$Q(\beta_{m_1}) \geq (1 - \delta) e^{-t(R + \varepsilon_1)}.$$

Установлено, таким образом, что для любого $\varepsilon_1 > 0$ и $\delta > 0$ при достаточно большой длине блока t существует множество α элементов m и множества β_m элементов Z , определяемые для каждого m следующими тремя свойствами:

- 1) $Pr(\alpha) \geq 1 - \delta$,
- 2) $d(Z, m) \leq d^* + \varepsilon_1$, если $Z \in \beta_m$,
- 3) $Q(\beta_m) \geq (1 - \delta) e^{-t(R + \varepsilon_1)}$, если $m \in \alpha$.

Отсюда, очевидно, вытекает, что для любого $\varepsilon > 0$ и достаточно большого t будут справедливы неравенства

- 1) $Pr(\alpha) \geq 1 - \varepsilon$,
- 2) $d(Z, m) \leq d^* + \varepsilon$, если $Z \in \beta_m$.
- 3) $Q(\beta_m) \geq e^{-t(R + \varepsilon)}$,

так как можно выбрать ε_1 и δ достаточно малыми, чтобы удовлетворить этим упрощенным условиям, в которых используется одно и то же ε . Этим завершается доказательство леммы.

Прежде чем перейти к общей проблеме кодирования, рассмотрим систему передачи, схематически изображенную на рис. 4. Имеем

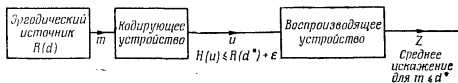


Рис. 4.

эргодический источник и задана мера искажения отдельных букв, которой соответствует функция $R(d)$. Требуется кодировать сообщения в последовательности u таким образом, чтобы исходные сообщения могли восстанавливаться воспроизводящим устройством со средним искажением, не превосходящим d^* (где d^* — некоторый фиксированный допустимый уровень искажения). Рассмотрим два устройства для блочного кодирования и для воспроизведения. На вход кодирующего устройства поступает вырабатываемая источником последовательность блоков длины t , а на выходе появляются блоки из алфавита u , соответствующие каждому возможному m -блоку.

Цель, которая поставлена здесь, состоит в том, чтобы кодировать сообщения так, чтобы при условии воспроизведения с искажением, не превосходящим d^* , сделать энтропию последовательностей u возможно более низкой. Энтропия, о которой здесь идет речь, есть энтропия на букву исходного сообщения или, иначе говоря, ее можно рассматривать как энтропию последовательности u в секунду, если считать, что источник создает одну букву в секунду.

Покажем, что для любого d^* и любого $\varepsilon > 0$ могут быть найдены такие кодирующие и воспроизводящие устройства, что

$H(u) \leq R(d^*) + \epsilon$, но при $\epsilon \rightarrow 0$ длина кодового блока, вообще говоря, возрастает. Этот вывод, конечно, тесно связан с нашей интерпретацией $R(d^*)$ как эквивалентной скорости источника при искажении d^* и легко вытекает из следующей теоремы.

Теорема 2. Пусть заданы эргодический источник, мера искажения d_{ij} и скорость $R(d)$, зависящая от искажения d (d — вычисленная на основе заданных частот отдельных букв); заданы $d^* \geq d_{\min}$ и $\delta > 0$. Тогда для любого достаточно большого t существует множество Λ , содержащее M слов длины t из алфавита Z со следующими свойствами:

$$1) \frac{1}{t} \log M \leq R(d^*) + \delta,$$

2) среднее искажение между t -словом длины t и ближайшим наименее искаженным к нему словом из множества Λ меньше или равно $d^* + \delta$.

Из этой теоремы следуют (если не учитывать в свойстве 2 δ , которое будет в дальнейшем устранено) изложенные выше выводы. А именно если в качестве кодирующего устройства использовать устройство, которое отображает любое t -слово на ближайший к нему элемент из множества Λ , а воспроизводящее устройство осуществляет просто тождественное преобразование, энтропия кодированной последовательности, приходящаяся на букву источника, не может превзойти $R(d^*) + \delta$, так как она имеет максимальное значение, равное $\frac{1}{t} \log M$, когда все M элементов из множества Λ равновероятны, а в силу условий теоремы

$$\frac{1}{t} \log M \leq R(d^*) + \delta.$$

Докажем эту теорему, используя метод случайного кодирования. Рассмотрим ансамбль способов выбора элементов множества Λ . Оценивая среднее искажение для этого ансамбля, найдем, что в нем существует по крайней мере один код с требуемыми свойствами.

Ансамбль кодов определен следующим образом. Для данного d^* существует множество переходных вероятностей $q_i(j)$, определяющих минимум R , равный $R(d^*)$. Множество вероятностей букв вместе с этими переходными вероятностями задает меру $Q(Z)$ в пространстве воспроизведенных слов. Мера Q для отдельной буквы алфавита Z , скажем буквы j , есть $\sum_i P_i q_i(j)$. Мера Q для слова, состоя-

щего из букв $j_1, j_2, \dots, j_t$, равна $Q(Z) = \prod_{k=1}^t (\sum_i P_i q_i(j_k))$.

В ансамбле кодов с кодовыми словами длины t соотнесем всевозможными способами слова длины t целым числам от 1 до M . Каждому целому числу соотнесено некоторое слово, скажем Z_1 ,

с вероятностью $Q(Z_i)$, причем вероятности для различных целых чисел статистически независимы. Это в точности такой же прием, как и при построении ансамбля случайных кодов для канала без памяти, за исключением того, что здесь целые числа отображаются на пространство Z с помощью меры $Q(Z)$. Таким образом, получаем множество кодов (если в алфавите Z имеются f букв, то в ансамбле будет f^M различных кодов), каждый из которых будет иметь соответствующую вероятность. Код, в котором целому числу i соотнесено Z_i , имеет вероятность $\prod_{i=1}^M Q(Z_i)$.

Теперь используем лемму 1, чтобы оценить среднее искажение для этого ансамбля кодов (используя для вычисления среднего искажения вероятности различных кодов). Заметим, во-первых, что если $Q(\beta)$ есть мера Q множества β слов Z в ансамбле кодов, то вероятность того, что это множество β не содержит кодовых слов, будет равна $[1 - Q(\beta)]^M$. Последнее выражение есть произведение вероятностей того, что множество β не содержит кодового слова 1, кодового слова 2 и т. д. Поэтому вероятность того, что множество β содержит по крайней мере одно кодовое слово, равна $1 - [1 - Q(\beta)]^M$. Теперь, обращаясь к лемме 1, видим, что для среднего искажения имеет место оценка

$$\bar{d} \leq \varepsilon d_{\max} + [1 - Q(\beta_m)]^M d_{\max} + (d^* + \varepsilon).$$

Здесь $d_{\max}$ — наибольшее искажение, возможное между буквами алфавитов m и Z . Первый член $\varepsilon d_{\max}$ определяется m -словами, которые не содержатся в множестве α . Их общая вероятность меньше или равна ε и при наличии их в сообщении среднее искажение меньше или равно $d_{\max}$. Второй член ограничивает сверху вклад, который обусловлен случаями, когда множество β_m , соответствующее сообщению m , не содержит ни одного кодового слова. Вероятность таких кодов в ансамбле не превосходит $[1 - Q(\beta_m)]^M$, а искажение не превосходит $d_{\max}$. Наконец, если сообщение содержится во множестве α и в β_m существует по крайней мере одно кодовое слово, то в соответствии с леммой 1 искажение ограничено величиной $d^* + \varepsilon$. Далее $Q(\beta_m) \geq e^{-t[R(d^*) + \varepsilon]}$ и для x такого, что $0 < x \leq 1$, справедливо соотношение

$$(1-x)^{\frac{1}{x}} = e^{\frac{1}{x} \log(1-x)} \leq e^{\frac{1}{x} \left(-x + \frac{x^2}{2}\right)} = e^{-1 + \frac{x}{2}} \leq e^{-\frac{1}{2}},$$

то (используя знакпеременность и монотонное убывание членов разложения логарифма) получаем

$$\begin{aligned} [1 - Q(\beta_m)]^M &\leq (1 - e^{-t[R(d^*) + \varepsilon]})^M = \\ &= [1 - e^{-t[R(d^*) + \varepsilon]}] e^{t[R(d^*) + \varepsilon]} e^{-t[R(d^*) + \varepsilon]} M \leq \\ &\leq e^{-\frac{M}{2}} e^{-t[R(d^*) + \varepsilon]}. \end{aligned}$$

Если взять M (число точек) равным по величине $e^{t[R(d^*)+2\varepsilon]}$ (или, если это не целое число, равным наименьшему целому, превосходящему эту величину), то выражение, приведенное выше, будет равно $e^{-\frac{1}{2}e^{t\varepsilon}}$. Следовательно, при таком выборе M получаем оценку для среднего искажения

$$\bar{d} \leq \varepsilon d_{\max} + e^{-\frac{1}{2}e^{t\varepsilon}} d_{\max} + d^* + \varepsilon \leq d^* + \delta$$

при условии, что в соответствии с леммой 1 ε будет выбрано достаточно малым для того, чтобы удовлетворить неравенству $(\varepsilon d_{\max} + 1) \leq \delta/2$, а t — достаточно большим, чтобы удовлетворить нера-

венству $e^{-\frac{1}{2}e^{t\varepsilon}} d_{\max} \leq \delta/2$. Кроме того, ε и t должны быть выбраны так, чтобы число M , наименьшее из целых чисел, большее или равное $e^{t[R(d^*)+2\varepsilon]}$, было бы меньше или равно $e^{t[R(d^*)+\delta]}$. Поскольку лемма 1 справедлива для всех достаточно больших t и любого положительного ε , все эти требования могут быть удовлетворены одновременно.

Таким образом, утверждения теоремы доказаны для среднего искажения ансамбля кодов. Отсюда следует, что существует по крайней мере один специфический код в этом ансамбле со средним искажением, ограниченным $d^* + \varepsilon$. Этим завершается доказательство теоремы.

С л е д с т в и е. Теорема 2 остается справедливой, если в утверждении 1) δ заменить на 0. Она также остается справедливой, если в утверждении 1) сохранить δ , а в утверждении 2) заменить его на 0, при условии что в этом случае $d^* > d_{\min}$, т. е. наименьшего d , для которого определена функция $R(d)$.

Этим следствием устанавливается, что можно достичь или даже превзойти одну координату точки кривой $R(d)$ и аппроксимировать сколь угодно точно вторую координату, за исключением, быть может, точки с абсциссой $d_{\min}$. Чтобы доказать это первое утверждение следствия, заметим прежде всего, что оно справедливо для $d^* \geq d_1$, такого, что $R(d_1) = 0$. Действительно, можно достигнуть точки $\bar{d} = d_1$ при $M = 1$ и длине кодового слова 1, образованного из такой буквы алфавита Z , которая задает среднее искажение для этого кода, не превосходящее d_1 . При $d_{\min} \leq d^* < d_1$ применим теорему 2 для аппроксимации $d^{**} = d^* + \delta/2$. Поскольку функция $R(d)$ строго убывающая, то эта аппроксимация приведет к кодам с $\bar{d} \leq d^* + \delta$ и $\frac{1}{t} \log M \leq R(d^*)$, если только δ в теореме 2 сделать достаточно малым.

Подобным же образом получается второе упрощение, указанное в следствии. Выбираем d^{**} , несколько меньшее, чем желаемое d^* ,

т. е. такое, что $R(d^{**}) = R(d^*) + \frac{\delta}{2}$, и используем теорему 2 для аппроксимации этой точки кривой.

Предположим теперь, что имеется канал без памяти с пропускной способностью C . В соответствии с теоремой кодирования для таких каналов возможно создать коды и системы декодирования, позволяющие передавать со скоростью на одну букву, приближающейся к C , и вероятностью ошибки, меньшей или равной ϵ_1 , для любого $\epsilon_1 > 0$. Комбинируя такой код для канала с кодом указанного выше типа для источника с заданным уровнем искажения d^* , получим следующий результат.

Теорема 3. *Даны источник, характеризуемый функцией $R(d)$, и канал без памяти с пропускной способностью $C > 0$; даны $\epsilon > 0$ и $d^* > d_{\min}$. Тогда существует для достаточно больших t и n блочный код, который соотносит слова длины t источника словам длины n на входе канала, и декодирующее устройство, которое отображает слова на выходе канала на воспроизводимые слова длины t . При этом удовлетворяются условия*

- 1) $\bar{d} \leq d^*$,
- 2) $\frac{nC}{t} \leq R(d^*) + \epsilon$.

Таким образом, можно достигнуть желаемого уровня искажения d^* (большого, чем $d_{\min}$) и в то же время приблизиться к использованию канала для передачи со скоростью, соответствующей $R(d^*)$. Это можно сделать, подобно тому как это сделано в следствии к теореме 2, аппроксимируя кривую $R(d)$ в точке левее d^* , скажем, в $R(d^* - \delta)$. Такой код будет иметь $M = e^{t[R(d^* - \delta) + \delta_1]}$ слов, где δ_1 может быть сделано как угодно малым при помощи выбора достаточно большого t . Код для канала образуется из M слов длины n , где n — наибольшее целое число, удовлетворяющее условию $\frac{nC}{t} \leq R(d^* - \delta) + \delta_1$. Выбирая t достаточно большим, можно сделать вероятность ошибки как угодно близкой к нулю, так как это соответствует скорости передачи, меньшей пропускной способности канала. Комбинация этих двух кодов приводит к общему коду со средним искажением, не превышающим d^* .

Двойственность источника и канала

Существует любопытная и поучительная двойственность между свойствами источника с мерой искажения и свойствами канала. Эта двойственность проявляется в большей степени, если рассматривается канал, в котором задана «цена», связанная с различными входными

буквами, и стоит задача нахождения пропускной способности при условии, что ожидаемая цена не превысит определенной величины. Пусть, например, входная буква i имеет цену α_i , и наша задача сводится к нахождению пропускной способности при дополнительном ограничении $\sum_i P_i \alpha_i \leq \alpha$, где P_i — вероятность появления на входе

буквы i . Эта проблема математически равносильна проблеме максимизации взаимной информации при вариации P_i с линейным неравенством в качестве ограничения. Решение этой проблемы дает пропускную способность канала как функцию цены $C(\alpha)$. Легко показать, что эта функция *выпукла вверх*. Решение этой проблемы соответствует в определенном смысле нахождению источника, согласованного с каналами и с заданной «ценой».

С другой стороны, вычисление функции $R(d)$ для источника математически равносильно минимизации взаимной информации изменением $q_i(j)$ опять-таки при линейном неравенстве в качестве ограничения. В результате получаем выпуклую вниз функцию $R(d)$. Решение этой проблемы соответствует нахождению канала, согласованного с источником и заданным уровнем искажения. Эта двойственность может быть прослежена и далее, и она относится к двойственности между прошлым и будущим и понятиями управления и знания. Так, можно обладать знаниями о прошлом, но нельзя управлять им. Можно управлять будущим, не зная его.

Численные расчеты для некоторых простых каналов

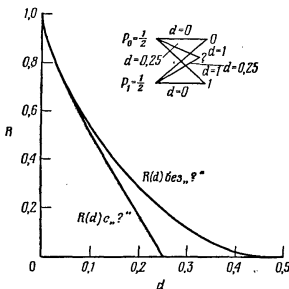
В этом разделе будут приведены некоторые численные расчеты для ряда простых каналов и источников. Рассмотрим сначала двоичный источник с независимыми и равновероятными буквами и предположим, что мерой искажения является вероятность ошибки (на знак). Такой источник попадает в класс источников, для которого можно дать простое и явное решение. Действительно, кривая $R(d)$ выражается в этом случае формулой

$$R(d) = 1 + d \log_2 d + (1 - d) \log_2 (1 - d).$$

Ясно, что это пропускная способность симметричного двоичного канала с вероятностями d и $(1 - d)$, поскольку такие значения $q_i(j)$ решают проблему минимизации.

Соответствующий график $R(d)$ показан на рис. 5. Там же построен ряд точек, соответствующих особенно простым кодам для двоичного канала в предположении, что шум в канале отсутствует. Расположение точек позволяет высказать некоторые соображения о том, насколько хорошо можно аппроксимировать нижнюю границу сравнительно простыми средствами. Например, точка, где искажение $d = 0$, получается при скорости $R = 1$ просто посылкой по кана-

нами блоков 3, 7, 15 и 31. Эти коды используются здесь в обратном порядке — любое сообщение в 15-мерном кубе, например, передается по каналу как сообщение из одиннадцати знаков, образующих ближайшую кодовую точку. На приемном конце соответствующее пятнадцатизначное сообщение восстанавливается. Оно может отличаться от исходного сообщения самое большее в одной позиции.



Р и с. 7.

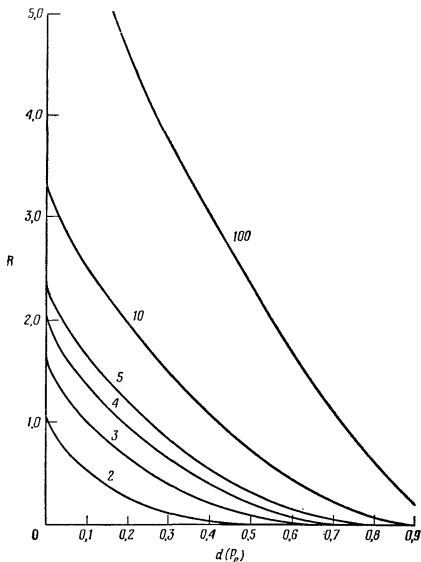
Следовательно, в этом случае отношение числа букв сигнала к числу букв сообщения равно 11/15, и легко найти, что вероятность ошибки будет 1/16. Эта серия точек дает близкую аппроксимацию нижней границы.

Промежутки между точками этих дискретных последовательностей можно заполнить при помощи использования смешанных кодов. Например, можно поочередно использовать два кода или, обобщая, перемешивать их в пропорции λ и $1 - \lambda$, где λ — любое рациональное число. Такое перемешивание приводит к коду с новым отношением R числа букв сообщения и сигнала, определяемым из соотношения $1/R = \lambda/R_1 + (1 - \lambda)/R_2$, где R_1 и R_2 — отношения для данных кодов, и с новой вероятностью ошибки

$$P_e = \frac{\lambda R_1 P_{e1} + (1 - \lambda) R_2 P_{e2}}{\lambda R_1 + (1 - \lambda) R_2}.$$

Эта интерполяция определяет выпуклую кривую между любыми двумя кодовыми точками. Будучи применена к отмеченным на рис. 5 ряду простых кодов и кодов, корректирующих одиночную ошибку, она дает интерполяционные линии, изображенные пунктиром.

В этой связи рассматривался также другой канал, а именно двоичный симметричный канал с пропускной способностью $C = 1/2$, для которого вероятность того, что знак принимается правильно,



Р и с. 8.

равна 0,89, и неправильно — 0,11. В этом случае совокупность точек, соответствующих простым кодам (рис. 6), действительно соприкасается с нижней границей в точке $R = 1/2$, так как сам канал без кодирования обуславливает как раз такую вероятность ошибки. Любой симметричный двоичный канал имеет одну точку, которая может быть в точности достигнута путем непосредственной передачи.

На рис. 7 показана кривая $R(d)$ для другого простого случая: источник является двоичным с независимыми буквами, а воспроизводимый алфавит Z состоит из 3-х букв «0», «1» и «?». Мера искажения равна 0 для правильного знака, 1 для неправильного знака и 0,25 для символа «?». На том же рисунке для сравнения показана кривая $R(d)$ для случая, когда символ «?» на выходе невозможен.

На рис. 8 приведены кривые для источников независимых букв с различным числом равновероятных букв в алфавите (2, 3, 4, 5, 10, 100). Здесь снова в качестве меры искажения принята вероятность ошибки (на знак). Если число букв в алфавите равно b , функция $R(d, b)$ выражается формулой

$$R(d, b) = \log_2 b + d \log_2 d + (1-d) \log_2 \frac{1-d}{b-1}.$$

Обобщение на случаи непрерывных алфавитов

Дадим теперь краткий набросок обобщения понятия меры искажения отдельной буквы на случаи, в которых входной и выходной алфавиты не ограничиваются конечными множествами, а изменяются в произвольных пространствах.

Пусть $A = \{m\}$ — алфавит сообщения и $B = \{z\}$ — алфавит воспроизводимых букв. Для каждой пары (m, z) из этих алфавитов задано неотрицательное число $d(m, z)$ — искажение, возникающее, когда m воспроизводится как z . Далее предположим, что над борелевским полем подмножеств пространства A определена вероятностная мера P . Наконец, потребуем, чтобы для каждого z , принадлежащего множеству B , $d(m, z)$ была измеримой функцией с конечным математическим ожиданием.

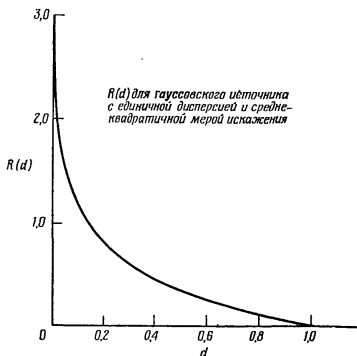
Рассмотрим конечный набор точек $z_i (i = 1, \dots, l)$ из пространства B и зададим измеримые переходные вероятности $q(z_i|m)$ (т. е. для каждого i функция $q(z_i|m)$ есть измеримая функция в пространстве A). При таком выборе z_i и задании $q(z_i|m)$ взаимная информация и среднее искажение определяются следующим образом:

$$R = \sum_i \int q(z_i|m) \log \frac{q(z_i|m)}{\int q(z_i|m) dP(m)} dP(m),$$

$$d = \sum_i \int d(m, z_i) q(z_i|m) dP(m).$$

Определим скорость $R(d^*)$ при заданном d^* для такого случая как наибольшую нижнюю грань R , когда варьируется совокупность точек z_i (как по выбору точек, так и по их числу), а функция $q(z_i|m)$ варьируется по измеримым переходным вероятностям; при этом уровень искажения не превышает d^* .

Большая часть выводов, полученных для случая конечного алфавита, легко проверяются и при этом обобщении. В частности, остается справедливым свойство выпуклости кривой $R(d)$. Действительно, пусть значение $R(d)$ может быть аппроксимировано



Р и с. 9.

с точностью ε с помощью выбора z_i и $q(z_i|m)$, а значение $R(d')$ с помощью выбора z'_i и $q'(z'_i|m)$. Рассмотрим совокупность точек z''_i , образованную из объединения точек z_i и z'_i , и переходные вероятности $q''(z''_i|m) = \frac{1}{2} [q(z_i|m) + q'(z'_i|m)]$ (заменяя $q(z''|m)$ или $q'(z''|m)$ нулем в точках z'' , где они не определены). Благодаря выпуклости и линейности d такой выбор приводит к точке с координатами $d'' = \frac{1}{2}d + \frac{1}{2}d'$ и R'' , находящейся в ε -окрестности средней точки отрезка прямой, соединяющего точку $d, R(d)$ и точку $d', R(d')$, или ниже ее. Поскольку ε может быть сделано произвольно малым, наибольшая нижняя грань $R(d'')$ находится в этой средней точке или ниже ее.

В общем случае, однако, кривая $R(d)$ не всегда имеет конечный предел при d , убывающем до своего наименьшего возможного значения. Кривая эта может, например, иметь вид, указанный на рис. 9, где $R(d)$ стремится к бесконечности по мере того, как d стремится

к $d_{\min}$. С другой стороны, при сформулированных условиях существует конечное значение $d_{\max}$, для которого $R(d_{\max}) = 0$. Это значение d дается выражением

$$d_{\max} = \liminf_z E[d(m, z)].$$

Обратная часть теоремы кодирования доказывается по существу точно так же, как и в случае конечного алфавита. Предполагается только, что допустимым кодирующим функциям, отображающим сообщения на входные сигналы, соответствуют измеримые подмножества пространства сообщений. (Если не сделать этого предположения, то вообще невозможно даже определить понятие «среднее искажение».) Замена в соответствующих неравенствах сумм, определенных в пространстве A на интегралы, приводит к доказательству соответствующей обратной теоремы.

Те же рассуждения, что и в случае конечного алфавита, могут быть использованы здесь и для доказательства прямой теоремы кодирования с дополнительным введением ε для аппроксимации нижней грани $R(d)$ при конечном наборе точек z_i . Выбирается множество точек z_i , которые аппроксимируют кривую $R(d)$ с точностью до ε , а затем применяется метод случайного кодирования. Единственный вопрос, который должен быть исследован несколько отличным методом, связан с членом $d_{\max}$. К каждому коду в ансамбле кодов можно добавить некоторую точку, например z_0 , и заменить $d_{\max}$ на $E[d(m, z_0)]$, являющееся конечной величиной¹⁾. Отсюда следуют все выводы теоремы.

Разностная мера искажения

Теперь рассмотрим специальный класс мер искажения для некоторых непрерывных случаев, имеющих важное значение, для которых могут быть получены более явные результаты. Пусть пространствами m и z будут теперь множества всех действительных чисел. Мера искажения $d(m, z)$ будет называться *разностной мерой искажения*, если она является функцией только разности $m - z$, т. е. $d(m, z) = d(m - z)$. Широко распространенными примерами такой меры могут служить критерий квадратичной ошибки $d(m, z) = (m - z)^2$ или критерий абсолютной ошибки $d(m, z) = |m - z|$.

Определим нижнюю грань $R(d)$ для разностной меры искажения. Прежде всего определим следующим образом функцию $\Phi(d)$ для данной разностной меры $d(u)$. Рассмотрим произвольную функцию распределения $G(u)$; пусть H — энтропия этого рас-

¹⁾ Существование z_0 , такой, что $E[d(m, z_0)] < \infty$ — это, конечно, особое дополнительное ограничение. — Прим. ред.

пределения, а d — среднее искажение между случайной величиной с данным распределением и нулем. Таким образом имеем:

$$H = - \int_{-\infty}^{+\infty} \log dG(u) dG(u),$$

$$d = \int_{-\infty}^{+\infty} d(u) dG(u).$$

Требуется найти максимум H варьированием по $G(u)$ при условии $d \leq d^*$. Верхняя грань, если она конечна, очевидно, достигается при некотором распределении и является максимумом. Этот максимум для заданного d^* обозначаем $\Phi(d^*)$, а соответствующую функцию распределения назовем максимизирующим распределением для этого d^* .

Теперь предположим, что имеется функция распределения (обобщенные вероятности букв) $P(m)$ в пространстве m с энтропией $H(m)$. Требуется показать, что

$$R(d) \geq H(m) - \Phi(d).$$

Пусть z_i — множество точек z и $q(z_i|m)$ — заданные переходные вероятности. Тогда взаимная информация между m и z может быть записана в виде

$$R = H(m) - \sum_i Q_i H(m|z_i),$$

где Q_i — полная вероятность z_i . Если d_i — среднее искажение между m и z_i , то

$$H(m|z_i) \leq \Phi(d_i).$$

Это вызвано тем, что $\Phi(d)$ является максимумом H при данном среднем искажении и что искажение есть функция только разности между m и z . Поэтому $\Phi(d_i)$ является максимальным значением $H(m|z_i)$ при любом z_i . Следовательно,

$$R \geq H(m) - \sum_i Q_i \Phi(d_i).$$

Поскольку энтропия, рассматриваемая как функционал, определенный в пространстве функцией распределения, обладает свойством выпуклости, а d в том же пространстве линейно, то с помощью рассуждений, аналогичных приведенным выше, можно доказать, что $\Phi(d)$ является выпуклой вниз функцией. Из этого вытекает, что $\sum_i Q_i \Phi(d_i) \leq \Phi(\sum_i Q_i d_i) = \Phi(d)$, где d — среднее искажение при выбранных z_i и заданных переходных вероятностях.

Отсюда следует, что

$$R \geq H(m) - \Phi(d).$$

Поскольку это справедливо для любого выбора z_i и $q(z_i|m)$, наше утверждение доказано.

Если для некоторых значений $P(m)$ и $d(u)$ допустимы такие выборы z_i и $q(z_i|m)$, при которых величина R как угодно мало отличается от правой части последнего неравенства, то ясно тогда, что эта правая часть есть $R(d)$. Таким, например, является случай, когда распределение $P(m)$ гауссовское и $d(u) = u^2$ (т. е. мерой искажения является средняя квадратичная ошибка). Предположим, что сообщение имеет дисперсию σ^2 и рассмотрим в пространстве z гауссовское распределение с нулевым средним и дисперсией, равной $\sigma^2 - d$. (Ясно, что если $\sigma^2 - d \leq 0$, то $R(d) = 0$, так как в качестве воспроизводимого сообщения может быть использована единственная точка z , равная нулю.) Пусть условное распределение вероятностей $q(m|z)$ — гауссовское с дисперсией d . Это предположение согласуется с гауссовским характером распределения $P(m)$, так как свертка нормальных распределений дает нормальное распределение с дисперсией, равной сумме дисперсий. Тогда условная вероятностная мера $q(z|m)$ также нормальна.

Простой расчет показывает, что при этом величина R достигает нижней грани, указанной выше. Получающаяся при этом кривая $R(d)$ имеет вид

$$R(d) = \begin{cases} \log \frac{\sigma}{\sqrt{d}}, & d \leq \sigma^2, \\ 0, & d > \sigma^2. \end{cases}$$

Для $\sigma^2 = 1$ она показана на рис. 9.

Определение меры локального искажения

До сих пор рассматривалась мера искажения d_{ij} [или $d(m, z)$], которая зависит только от сравнения буквы сообщения и соответствующей воспроизводимой буквы. Это побуквенное искажение усреднялось по длине сообщения и по множествам возможных исходных и воспроизводимых сообщений. Однако во многих практических случаях такая мера не является достаточно общей. Значимость одного и того же типа ошибки часто зависит от контекста.

Например, при передаче цитаты биржевого бюллетеня, скажем, «А.Т. & Т. 5900 акций отдать по 194», ошибка в цифре 9 в числе акций 5900 обычно значительно менее существенна, чем ошибка в цифре 9 в цене 194.

Рассмотрим теперь меру искажения, которая зависит от локального контекста и сравним блоки из g букв сообщения с соответствующими блоками из g букв воспроизводимого сообщения.

Мера локального искажения периода g есть неотрицательная функция $d(m_1, m_2, \dots, m_g; z_1, z_2, \dots, z_g)$ последовательностей букв сообщений длины g и последовательностей букв воспроизводимых сообщений длины g (возможно, из алфавита, отличного от алфавита исходных сообщений). Искажение между $m = m_1, m_2, \dots, m_t$ и $z = z_1, z_2, \dots, z_t$ ($t \geq g$) определяется выражением

$$d(m, Z) = \frac{1}{t-g+1} \sum_{h=1}^{t-g+1} d(m_h, m_{h+1}, \dots, m_{h+g-1}; z_h, z_{h+1}, \dots, z_{h+g-1}).$$

Искажение блокового кода, в котором сообщение m и его воспроизводимый вариант Z появляются с вероятностью $P(m, Z)$, определяется как

$$d = \sum_{m, Z} P(m, Z) d(m, Z).$$

Другими словами, считается, что оценка системы в целом, когда задана мера локального искажения, получается путем усреднения искажений для всех сравниваемых блоков длины g с их вероятностями как весовыми множителями.

Функции $R_n(d)$ и $R(d)$ для меры локального искажения и эргодического источника

Предположим, что заданы эргодический источник сообщения и мера локального искажения. Рассмотрим n -буквенные блоки сообщений с их вероятностями (определяемыми источником) вместе с возможными блоками Z воспроизводимых сообщений длины n . Пусть значения переходных вероятностей от m -блоков к Z -блокам заданы как $q(Z|m)$. При таком задании можно вычислить две величины: 1) среднюю взаимную информацию на букву

$$R = \frac{1}{n} E \left(\log \frac{q(Z|m)}{Q(Z)} \right)$$

и 2) среднее искажение

$$d = \sum_{m, Z} P(m, Z) d(m, Z).$$

Варьируя $q(Z|m)$ при условии $d \leq d^*$, можно в принципе найти минимум величины R для каждого значения d^* . Назовем эту величину $R_n(d^*)$.

Возникающая здесь проблема минимизации совпадает с исследованной ранее, если рассматривать m и Z как отдельные буквы

(большого) алфавита. Поэтому различные выводы, полученные ранее, могут быть использованы и здесь. В частности, $R_n(d)$ есть выпуклая вниз функция.

Определим теперь для данного источника функцию скорости искажения относительно рассматриваемой меры искажения

$$R(d) = \liminf_{n \rightarrow \infty} R_n(d).$$

Можно показать прямыми, но трудоемкими рассуждениями, которые не приводятся, что в этом равенстве можно заменить нижний предел на обычный. Другими словами, $R_n(d)$ стремится к пределу, когда $n \rightarrow \infty$.

Теперь можно доказать теорему кодирования для произвольного эргодического источника с мерой локального искажения.

Прямая теорема кодирования для меры локального искажения

Теорема 4. *Предположим, что даны эргодический источник, мера локального искажения и функция $R(d)$. Пусть K — двоичный канал без памяти с пропускной способностью C , d^* — величина искажения и пусть ε — положительное число. Тогда существует блок-код с искажением, меньшим или равным $d^* + \varepsilon$, и скоростью передачи, не меньшей $(C/R - \varepsilon)$ букв сообщения на букву в канале.*

Доказательство. Выберем n_1 так, что $R_{n_1}(d^*) - R(d^*) < \frac{\varepsilon}{3}$ и $\frac{g}{n_1} d_{\max} < \frac{\varepsilon}{3}$. Теперь рассмотрим блоки длины n_1 в качестве «букв» расширенного алфавита. Применяя теорему 3, построим блок-код, использующий достаточно длинные последовательности этих букв, передаваемых со скоростью, близкой (скажем, с точностью $\varepsilon/3$) к $R_{n_1}(d^*)/C$ (на букву исходного сообщения) и с искажением, меньшим чем $d^* + \frac{\varepsilon}{3}$. Необходимо помнить, что это искажение основано на сравнении искажений отдельных букв. Однако искажение, определяемое локальной мерой, будет отличаться от «побуквенного искажения» тем, что, помимо искажений, обусловленных буквами сообщения, оно содержит еще дополнительные искажения, обусловленные блоками длины g , составленными из букв исходного алфавита, принадлежащих двум соседним буквам нового алфавита. Таких дополнительных членов будет g на каждые n_1 букв сообщения. Поэтому разница между искажениями, определяемыми локальной мерой и мерой искажения отдельной буквы, не превышает $\frac{g}{n_1} d_{\max} < \frac{\varepsilon}{3}$. Отсюда следует, что эти коды позволяют

передавать сообщения со скоростью, отличающейся не более чем на ε от $R(d^*)$, и с искажением, отличающимся не более чем на ε от значения d^* .

Обратная теорема кодирования

Теорема 5. Предположим, что даны эргодический источник, мера локального искажения и скорость $R(d)$, зависящая от искажения d . Пусть K — канал без памяти с пропускной способностью C , d^* — величина искажения и пусть ε — положительное число. Тогда существует такое t_0 , что при любой передаче $t \geq t_0$ букв сообщения с искажением, не превышающим d^* при n -кратном использовании канала, удовлетворяется неравенство

$$\frac{n}{t} C \geq R(d^*) - \varepsilon.$$

Это означает, что пропускная способность канала, измеряемая в битах на букву сообщения, при длительных передачах будет близкой к значению $R(d^*)$.

Доказательство. Выберем t_0 так, чтобы для $t \geq t_0$ было бы справедливо неравенство $R_t(d) \geq R(d) - \varepsilon$. Это возможно, поскольку значение $R(d)$ определено как нижний предел от $R_t(d)$. Предположим, что для такого $t \geq t_0$ имеется код, который отображает последовательности m , состоящие из t букв сообщения, в последовательности X из n букв сигнала на входе канала и отображает последовательности Y из n букв выходного сигнала в последовательности Z воспроизводимых сообщений. Переходные вероятности канала образуют функцию $P(Y|X)$. Кроме того, заданы кодирующие и декодирующие функции, определяемые как $X = f(m)$ и $Z = g(Y)$. Наконец, источник задает вероятности $P(m)$ последовательностей сообщений m . Посредством кодирующей функции $f(m)$ задается множество вероятностей $P(X)$ последовательностей на входе канала. Если пропускная способность канала равна C , то средняя взаимная информация $R(X, Y)$ между последовательностями на входе и выходе канала должна удовлетворять соотношению

$$R(X, Y) = E \log \frac{P(X|Y)}{P(X)} \leq nC,$$

так как nC — максимально возможное значение $R(X, Y)$ при варьировании $P(X)$. Поскольку X является функцией m , а Z — функция Y , то аналогично

$$R(m, Z) = E \log \frac{P(m|Z)}{P(m)} \leq R(X, Y) \leq nC.$$

Кодирование, о котором идет речь, по существу равнозначно множеству условных вероятностей перехода от последовательностей m к последовательностям Z , определяемым двумя кодирующими функциями и переходными вероятностями в канале. Если полное искажение меньше или равно d^* , то $tR_t(d^*) = \min_{P(Z|m)} R(m, Z)$

будет заведомо меньше или равно некоторому частному значению $R(m, Z)$, определяемому вероятностями, задаваемыми каналом и методом кодирования (наличие множителя t объясняется тем, что $R_t(d)$ измеряется относительно буквы сообщения, в то время как $R(m, Z)$ относится к последовательности длины t). Таким образом,

$$tR_t(d^*) \leq R(m, Z) \leq nC,$$

$$t(R(d^*) - \varepsilon) \leq nC,$$

$$\frac{n}{t} C \geq R(d^*) - \varepsilon.$$

Этим заканчивается доказательство теоремы.

Заметим, что, как видно из метода доказательства, используемый код опять-таки не обязательно должен быть блоковым кодом, нужно только, чтобы после n -кратного использования канала были записаны t воспроизведенных букв. Если имеется какой-либо неравномерный код, который непрерывно используется, начиная с момента нуль, то последнее неравенство теоремы будет иметь место для любого конечного момента времени, к которому воспроизведены t_0 букв сообщения и, конечно, при неограниченном удлинении сравниваемых блоков $\varepsilon \rightarrow 0$. Это утверждение можно обобщить даже на неравномерные коды, для которых после n -кратного использования канала число воспроизведенных букв сообщения есть случайная величина, зависящая, быть может, от специфики сообщения и специфики случайного состояния канала. Если существует, что обычно имеет место для таких кодов, средняя скорость передачи, при которой после n -кратного использования канала с вероятностью, близкой к единице, будут воспроизведены t букв, где t заключено в пределах $t_1(1 - \delta)$ и $t_1(1 + \delta)$ (при $\delta \rightarrow 0$, когда $n \rightarrow \infty$), то по существу можно применить ту же самую теорему, заменяя t на t_1 .

Канал с памятью

Наконец заметим, что, хотя выше рассматривался канал без памяти, можно получить весьма похожие прямые и обратные результаты для каналов с памятью.

Для канала с памятью можно определить пропускную способность C_n для первых n использований канала, если канал исходно находился в состоянии s_0 . Это значение C_n есть $1/n$ часть максимума

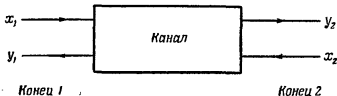
средней взаимной информации между последовательностями входного сигнала длины n и результирующими последовательностями выходного сигнала при варьировании вероятностей последовательностей входного сигнала длины n . Нижняя грань для значения P_e после n -кратного использования канала оценивается следующим образом:

$$P_e \geq \Phi\left(\frac{R}{C_n}\right).$$

Можно также определить пропускную способность C для такого канала, как $C = \lim_{n \rightarrow \infty} \sup C_n$. Тогда прямая часть теоремы утверждает, что можно найти блочный код произвольной длины, при котором скорость больше или равна R , а вероятность ошибки на знак меньше или равна $\Phi\left(\frac{R}{C}\right) + \varepsilon$. Конечно, в большинстве интересующих нас каналов влияние предьстории уменьшается так, что $C_n \rightarrow C$, когда $n \rightarrow \infty$. Для каналов без памяти $C_n = C$ для всех n .

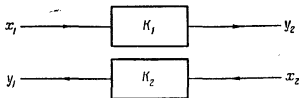
1. Введение

Двусторонний канал связи схематически изображен на рис. 1. Здесь x_1, y_1 являются соответственно входными и выходными буквами на первом конце; x_2, y_2 — входные и выходные буквы на втором



Р и с. 1.

конце канала. Пусть, скажем, раз в секунду могут быть выбраны по одной букве x_1 и x_2 из соответствующих алфавитов, и затем эти буквы поступают на входы канала; тогда на выходах канала могут

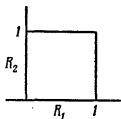


Р и с. 2.

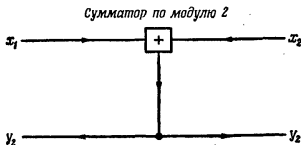
наблюдаться y_1 и y_2 , которые статистически связаны с входными буквами x_1 и x_2 , а также, возможно, и с другими предшествующими буквами на входах и выходах, если канал имеет память. Проблема состоит в том, чтобы вести передачу в обоих направлениях как можно эффективнее. В частности, желательно выяснить, какие пары скоростей передачи сигналов R_1 и R_2 для обоих направлений могут быть достигнуты при произвольно малой вероятности ошибки.

¹⁾ Shannon C., Two-way communication channels, Proc. 4-th Berkeley symposium on Math. Statist. and probability, 1960, I, 1962, 611.

Перед тем как ввести точные определения, рассмотрим несколько простых примеров. На рис. 2 показан двусторонний канал, распадающийся на два независимых односторонних двоичных канала без шума K_1 и K_2 . Здесь каждая из величин x_1 , x_2 , y_1 и y_2 может принимать лишь два значения и действие канала характеризуется соотношениями $y_2 = x_1$ и $y_1 = x_2$. Передачу можно вести со скоростью один бит в секунду. При этом можно найти такие коды, что соответствующие скорости (R_1, R_2) аппроксимируют сколь угодно точно



Р и с. 3.



Р и с. 4.

любую точку в квадрате рис. 3 и вероятность ошибки является произвольно малой (в данном случае нулевой). На рис. 4 входные и выходные сигналы снова могут принимать лишь по два значения и канал характеризуется соотношением $y_1 = y_2 = x_1 + x_2 \pmod{2}$. Здесь также можно вести передачу со скоростью один бит в секунду по обоим направлениям одновременно, но ее нужно вести несколько более утонченным методом. Хотя в качестве x_1 и x_2 могут быть переданы произвольные двоичные числа, но при декодировании наблюдаемые значения y должны быть прокорректированы, чтобы исключить влияние переданного значения x . Так, чтобы определить переданное значение x_2 надо к наблюдаемому значению y_1 прибавить точное значение переданного $x_1 \pmod{2}$. Конечно, и здесь можно получить пары скоростей, меньшие чем $(1,1)$, и снова аппроксимировать любую точку квадрата рис. 3.

В третьем примере входные буквы x_1 и x_2 принимают значения из *троичного* алфавита $(0,1,2)$, а выходные буквы y_1 и y_2 из двоичного (a,b) . Предположим, что условные вероятности различных пар (y_1, y_2) на выходах канала при условии, что заданы пары (x_1, x_2) на входах канала, принимают значения, приведенные в табл. 1. Легко видеть, что, используя на первом конце только $x_1 = 0$, можно вести передачу в направлении 2—1 со скоростью один бит в секунду, используя при этом только буквы 1 и 2 на втором конце, которые при этом перейдут соответственно в буквы a и b на первом конце канала. Аналогично если положить $x_2 = 0$, то в направлении 1—2 передача может вестись со скоростью один бит в секунду. Деля общее

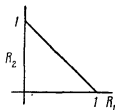
время передачи в отношении λ к $1 - \lambda$ для использования этих двух методов, можно вести передачу по обоим направлениям со средними скоростями $R_1 = 1 - \lambda$ и $R_2 = \lambda$ соответственно. Таким образом, можно найти коды, аппроксимирующие любую точку в треугольной

Т а б л и ц а 1

$x_1 x_2$		Выходные пары			
		aa	ab	ba	bb
Входные пары	00	$1/4$	$1/4$	$1/4$	$1/4$
	01	$1/2$	$1/2$	0	0
	02	0	0	$1/2$	$1/2$
	10	$1/2$	0	$1/2$	0
	11	$1/4$	$1/4$	$1/4$	$1/4$
	12	$1/4$	$1/4$	$1/4$	$1/4$
	20	0	$1/2$	0	$1/2$
	21	$1/4$	$1/4$	$1/4$	$1/4$
	22	$1/4$	$1/4$	$1/4$	$1/4$

области, изображенной на рис. 5. Нетрудно видеть, и это будет следовать из дальнейшего, что вне этого треугольника не имеется точек, которые могут быть аппроксимированы кодами с произвольно малой вероятностью ошибки.

Передача по двум направлениям в этом канале может быть названа несовместимой. Передача в прямом направлении (1—2



Р и с. 5.

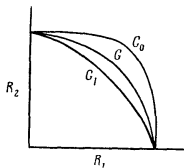
возможна, если только положить $x_2 = 0$. В других случаях значения всех букв y_2 полностью искажаются шумом. Наоборот, передача в обратном направлении (2—1) возможна, если только положить $x_1 = 0$. Ситуация, возникающая здесь, аналогична обычной реальной двусторонней системе связи: пара радиотелефонных станций с кнопками для перехода на передачу устроена так, что, когда эта кнопка нажата, прием выключен.

Четвертый простой пример двустороннего канала — предложенный Блекуэллом — двоичный умножающий канал. Здесь все буквы на входах и выходах являются двоичными величинами и канал характеризуется соотношением $y_1 = y_2 = x_1 x_2$. Область пар аппроксимируемых скоростей для этого канала точно не известна, но позднее для нее будут найдены некоторые границы.

В данной работе будут рассмотрены свойства кодирования для двусторонних каналов. В частности, будут найдены внутренние и внешние границы для области пар аппроксимируемых скоростей (R_1, R_2) и границы областей, точки которых могут быть аппроксимированы парами скоростей при нулевой вероятности ошибки. Будут рассмотрены некоторые топологические свойства этих границ и в заключение дан способ описания области пар аппроксимируемых скоростей в терминах некоторого предельного процесса.

2. Краткое изложение содержания

Сформулируем кратко и в несколько огрубленной форме основные результаты настоящей статьи. В ней показано, что для произвольного дискретного канала без памяти существует выпуклая область G аппроксимируемых скоростей, т. е. для любой точки



Р и с. 6.

(R_1, R_2) , лежащей внутри G , существуют коды, позволяющие вести передачу со скоростями, как угодно близкими к этой точке, и передача может быть осуществлена с произвольно малой вероятностью ошибки. Типичная форма такой области показана на рис. 6. Эта область ограничена средней кривой G и двумя отрезками осей. Кривая G может быть описана некоторым асимптотическим выражением, содержащим взаимную информацию длинных последовательностей входных и выходных букв.

Кроме того, найдем внутреннюю и внешнюю границы G_I и G_O , которые могут быть вычислены гораздо легче, так как они получаются просто процессом максимизации, относящимся к отдельным буквам в канале. G_O — это множество точек (R_{12}, R_{21}) , которые мож-

но получить, задавая произвольные совместные распределения $P\{x_1, x_2\}$ букв на входах канала и затем вычисляя по формулам:

$$\begin{aligned} R_{12} &= E \left(\log \frac{P\{x_1 | x_2, y_2\}}{P\{x_1 | x_2\}} \right) = \\ &= \sum_{x_1, x_2, y_2} P\{x_1, x_2, y_2\} \log \frac{P\{x_1 | x_2, y_2\}}{P\{x_1 | x_2\}}. \\ R_{21} &= E \left(\log \frac{P\{x_2 | x_1, y_1\}}{P\{x_2 | x_1\}} \right), \end{aligned} \quad (1)$$

где $E(\mu)$ — математическое ожидание величины μ . Внутренняя граница G_I может быть найдена аналогичным путем, но на совместные распределения накладываем условие независимости $P\{x_1, x_2\} = P\{x_1\}P\{x_2\}$. Тогда G_I есть выпуклая оболочка точек (R_{12}, R_{21}) , найденных при этом ограничении.

Будет показано, что в некоторых важных случаях эти границы совпадают, тогда область пропускной способности полностью определяется этими границами. Приводятся также некоторые примеры (двоичный умножающий канал), когда имеется различие между этими границами.

Наши три области G_I , G и G_0 все являются выпуклыми и отсекают одинаковые отрезки на осях. Эти отрезки представляют собой пропускные способности канала по каждому из двух направлений, если на противоположном входе фиксирована наилучшая буква (например, такое значение x_1 , которое максимизирует R_{21} при изменении $P\{x_2\}$). Для любой точки внутри G вероятность ошибки стремится к нулю с экспоненциальной скоростью вместе с ростом блоков длины n . Для любой внешней точки области G по крайней мере для одного из двух кодов вероятность ошибки будет ограничена снизу независимо от длины блоков.

Наконец, полученные результаты могут быть частично обобщены на некоторый класс каналов с памятью. Если существует внутреннее состояние канала, такое, что можно вернуться в это состояние через ограниченное число шагов (независимо от проведенной ранее передачи), то снова можно найти область пропускной способности G , обладающую аналогичными свойствами. Приводится предельное соотношение, определяющее эту область.

3. Основные определения

Двусторонний дискретный канал без памяти задается совокупностью переходных вероятностей $P\{y_1, y_2 | x_1, x_2\}$, где x_1, x_2, y_1 и y_2 принимают значения из соответствующих конечных алфавитов (алфавиты не обязательно одинаковые).

Пара блоковых кодов длины n такого канала для передачи M_1 сообщений в прямом направлении и M_2 в обратном состоит из двух множеств по n функций в каждом:

$$f_0(m_1), f_1(m_1, y_{11}), f_2(m_1, y_{11}, y_{12}), \dots, f_{n-1}(m_1, y_{11}, \dots, y_{1, n-1}), \quad (2)$$

$$g_0(m_2), g_1(m_2, y_{21}), g_2(m_2, y_{21}, y_{22}), \dots, g_{n-1}(m_2, y_{21}, \dots, y_{2, n-1}).$$

Здесь все функции f принимают значения в алфавите x_1 , все функции g — в алфавите x_2 , m_1 принимают значения от 1 до M_1 (сообщения в прямом направлении) и m_2 от 1 до M_2 (сообщения в обратном направлении). Наконец, y_{1i} (аналогично y_{2i}), $i = 1, 2, \dots, n-1$ принимают значения в алфавите y_1 (соответственно y_2). Функции f определяют, как должна быть выбрана следующая буква на первом конце канала, если известно сообщение m_1 , которое должно быть передано, и известны буквы на выходе первого конца канала $y_{11}, y_{12}, \dots$, полученные к этому моменту времени. Аналогично функции g по имеющейся информации в каждый момент процесса определяют, как должно кодироваться сообщение m_2 .

Система декодирования для пары блоковых кодов длины n состоит из пары функций $\varphi(m_1, y_{11}, y_{12}, \dots, y_{1n})$ и $\psi(m_2, y_{21}, y_{22}, \dots, y_{2n})$, которые принимают значения от 1 до M_2 и от 1 до M_1 соответственно.

Декодирующая функция φ представляет собой способ принятия решения о том, какое сообщение было передано со второго конца канала, на основе информации, имеющейся на первом конце канала. Заметим, что буквы передаваемой последовательности $x_{11}, x_{12}, \dots, x_{1n}$, хотя и известны на первом конце, не включены в качестве аргументов в декодирующую функцию, так как они полностью определяются, если знать кодирующие функции, сообщения m_1 и принятую последовательность $y_{11}, y_{12}, \dots, y_{1n}$.

Будем предполагать, если не оговорено противное, что все сообщения m_1 и m_2 являются независимыми и равновероятными с вероятностями соответственно $1/M_1$ и $1/M_2$. Также будем предполагать, что воздействия канала на последовательные буквы независимы, т. е.

$$\begin{aligned} P\{y_{11}, y_{12}, \dots, y_{1n}, y_{21}, y_{22}, \dots, y_{2n} | x_{11}, x_{12}, \dots, x_{1n}, x_{21}, x_{22}, \dots, x_{2n}\} = \\ = \prod_{i=1}^n P\{y_{1i}, y_{2i} | x_{1i}, x_{2i}\}. \end{aligned} \quad (3)$$

Это условие выражает отсутствие памяти у канала. Другими словами, это означает, что вероятность некоторого множества букв на обоих выходах канала при условии, что известны соответствующие

щие буквы на обоих входах канала, равна вероятности того же события при условии, что дополнительно известно любое количество предыдущих букв на входах канала.

Скоростями передачи сообщений, количество которых равно M_1 и M_2 по обоим направлениям в канале для пары блоковых кодов, называются величины R_1 и R_2 , определяемые так:

$$\begin{aligned} R_1 &= \frac{1}{n} \log M_1, \\ R_2 &= \frac{1}{n} \log M_2. \end{aligned} \quad (4)$$

Зная пару кодов, систему декодирования, условные вероятности, определяющие канал, и основываясь на наших предположениях относительно вероятности сообщений, в принципе возможно вычислить вероятность ошибки при данном коде. С этой целью можно сначала для каждой пары сообщений вычислить вероятности получения на выходах различных возможных последовательностей, если при передаче эти сообщения кодировались с помощью данных кодирующих функций. Применяя декодирующие функции, можно вычислить вероятность неверного декодирования. Осреднив последние вероятности по всевозможным сообщениям для каждого направления, можно получить, наконец, выражения для вероятностей ошибок P_{e1} и P_{e2} при передаче соответственно по обоим направлениям канала.

Будем говорить, что точка (R_1, R_2) принадлежит области пропускной способности G данного канала без памяти K , если для любого $\varepsilon > 0$ существует некоторая пара блоковых кодов и система декодирования, такие, что скорости передачи R_1^* и R_2^* удовлетворяют соотношению $|R_1 - R_1^*| < \varepsilon$ и $|R_2 - R_2^*| < \varepsilon$ и вероятности ошибок $P_{e1} < \varepsilon$ и $P_{e2} < \varepsilon$.

4. Средние взаимные скорости передачи информации

Двусторонний дискретный канал без памяти с конечными алфавитами был определен с помощью совокупности вероятностей перехода $P\{y_1, y_2 | x_1, x_2\}$, где x_1 и x_2 являются входными буквами соответственно первого и второго конца канала, а y_1 и y_2 соответствующие буквы на его выходах. Каждая из этих букв принимает значение из соответствующего конечного алфавита.

Если заданы (произвольные) статистически независимые распределения вероятностей $P\{x_1\}$ и $P\{x_2\}$ букв из алфавитов x_1 и x_2 , то можно определить соответствующие распределения вероятностей $P\{y_1\}$ и $P\{y_2\}$ для букв из алфавитов y_1 и y_2 . Действитель-

но, для четырех случайных величин x_1, x_2, y_1, y_2 имеем

$$P\{x_1, x_2, y_1, y_2\} = P\{x_1\}P\{x_2\}P\{y_1, y_2|x_1, x_2\}, \quad (5)$$

$$P\{y_1\} = \sum_{x_1, x_2, y_2} P\{x_1, x_2, y_1, y_2\}$$

и т. д.

Исходя из интуитивных соображений и аналогии с односторонними каналами, можно определить скорость передачи от первого конца к второму как $H(x_1) - H(x_1|x_2, y_2)$, т. е. разность между неопределенностью, или энтропией x_1 , и условной энтропией x_1 при условии, что известны величины, доступные на втором конце, а именно величины y_2 и x_2 . Таким образом, можно написать

$$\begin{aligned} R_{12} = H(x_1) - H(x_1|x_2, y_2) &= E \left[\log \frac{P\{x_1, x_2, y_2\}}{P\{x_1\}P\{x_2, y_2\}} \right] = \\ &= E \left[\log \frac{P\{x_1|x_2, y_2\}}{P\{x_1\}} \right]. \end{aligned} \quad (6)$$

$$\begin{aligned} R_{21} = H(x_2) - H(x_2|x_1, y_1) &= E \left[\log \frac{P\{x_1, x_2, y_1\}}{P\{x_2\}P\{x_1, y_1\}} \right] = \\ &= E \left[\log \frac{P\{x_2|x_1, y_1\}}{P\{x_2\}} \right]. \end{aligned} \quad (7)$$

Эти выражения при заданных распределениях вероятностей на входах канала представляют собой информацию между входом на одном конце и парой вход—выход на другом. Можно ожидать, что при соответствующем кодировании возможно вести передачу по обоим направлениям *одновременно* с произвольно малыми вероятностями ошибок со скоростями, как угодно близкими к R_{12} и R_{21} . Обобщая случай односторонних каналов, можно предположить, что построение таких кодов надо основывать на использовании распределений вероятностей $P\{x_1\}$ и $P\{x_2\}$ и, действительно, покажем, что, основываясь на вероятностях $P\{x_1\}$ и $P\{x_2\}$, можно найти коды, обладающие такими свойствами.

Однако область пропускной способности может быть больше, чем совокупность скоростей, которые можно получить таким способом. Грубо говоря, различие вытекает из того, что случайные величины x_1 и x_2 могут оказаться зависимыми. В этом случае соответствующие взаимные информации даются выражениями

$$H(x_2|x_1) - H(x_2|x_1, y_1) \text{ и } H(x_1|x_2) - H(x_1|x_2, y_2).$$

Приведенные выше формулы для R_{12} и R_{21} сводятся, конечно, к этим, когда x_1 и x_2 независимы.

5. Распределение информации

Применяемый здесь метод основан на случайном кодировании и методе, подобном тому, который развит в работе¹⁾ для одно-сторонних каналов. Рассмотрим n -кратное использование канала или, говоря математически, произведение вероятностных пространств. На вход поступают величины $X_1 = (x_{11}, x_{12}, \dots, x_{1n})$ и $X_2 = (x_{21}, x_{22}, \dots, x_{2n})$, а на выходе появляются $Y_1 = (y_{11}, y_{12}, \dots, y_{1n})$ и $Y_2 = (y_{21}, y_{22}, \dots, y_{2n})$, т. е. последовательности по n букв из соответствующих алфавитов.

Условные вероятности для этих блоков равны

$$P(Y_1, Y_2 | X_1, X_2) = \prod_k P(y_{1k}, y_{2k} | x_{1k}, x_{2k}). \quad (8)$$

Здесь использовалось предположение о том, что рассматривается канал без памяти, т. е. последовательные символы передаются независимо. Вероятностную меру на входных блоках X_1 и X_2 также определим как произведение мер, заданных на x_1 и x_2 . Таким образом,

$$P\{X_1\} = \prod_k P\{x_{1k}\},$$

$$P\{X_2\} = \prod_k P\{x_{2k}\}. \quad (9)$$

Отсюда уже следует, что другие вероятности также являются произведениями соответствующих вероятностей для отдельных букв. Так, например,

$$P\{X_1, X_2, Y_1, Y_2\} = \prod_k P\{x_{1k}, x_{2k}, y_{1k}, y_{2k}\},$$

$$P\{X_2 | X_1, Y_1\} = \prod_k P\{x_{2k} | x_{1k}, y_{1k}\}. \quad (10)$$

Взаимная (неусредненная) информация между X_1 и парой X_2, Y_2 может быть записана в виде суммы

$$\begin{aligned} I(X_1; X_2, Y_2) &= \log \frac{P\{X_1, X_2, Y_2\}}{P\{X_1\} P\{X_2, Y_2\}} = \log \frac{\prod_k P\{x_{1k}, x_{2k}, y_{2k}\}}{\prod_k P\{x_{1k}\} \prod_k P\{x_{2k}, y_{2k}\}} = \\ &= \sum_k \log \frac{P\{x_{1k}, x_{2k}, y_{2k}\}}{P\{x_{1k}\} P\{x_{2k}, y_{2k}\}}; \end{aligned} \quad (11)$$

$$I(X_1; X_2, Y_2) = \sum_k I(x_{1k}; x_{2k}, y_{2k}).$$

¹⁾ Shannon C., Certain results in coding theory for noisy channels. (Русский перевод см. стр. 509 данного сборника. — Прим ред.)

Таким образом, взаимная информация многомерных величин, как это обычно бывает в случае независимых величин, является суммой взаимных информаций одномерных величин. Так же, как обычно, можно считать, что взаимная информация является случайной величиной: $I(X_1; X_2, Y_2)$, которая принимает свои значения с вероятностями $P\{X_1, X_2, Y_2\}$. Функцию распределения для $I(X_1; X_2, Y_2)$ обозначим через $q_{12}(Z)$, а для $I(X_2; X_1, Y_1)$ — через $q_{21}(Z)$:

$$\begin{aligned} q_{12}(Z) &= P\{I(X_1; X_2, Y_2) \leq Z\}, \\ q_{21}(Z) &= P\{I(X_2; X_1, Y_1) \leq Z\}. \end{aligned} \quad (12)$$

Так как каждая из случайных величин $I(X_1; X_2, Y_2)$ и $I(X_2; X_1, Y_1)$ является суммой n независимых одинаково распределенных случайных величин, имеем обычные условия, когда можно применить различные варианты центральной предельной теоремы и закона больших чисел. Средние значения распределений q_{12} и q_{21} равны соответственно nR_{12} и nR_{21} , а дисперсии — соответствующим дисперсиям для одной буквы, умноженным на n . При $n \rightarrow \infty$ $q_{12}[n(R_{12} - \epsilon)] \rightarrow 0$ для любого фиксированного $\epsilon > 0$; аналогичное положение имеет место и для q_{21} . На самом деле, стремление к нулю является экспоненциальным по n , и определяется неравенством $q_{12}[n(R_{12} - \epsilon)] \leq \exp[-A(\epsilon)n]$.

6. Случайное кодирование для двусторонних каналов

После этих предварительных замечаний докажем существование кодов с некоторыми вероятностями ошибок, ограниченных выражениями, содержащими функции распределения q_{12} и q_{21} .

Будем строить некоторый ансамбль кодов, или более точно *пар кодов*: один код для передачи по направлению 1—2, другой — для 2—1. Установим границы для вероятностей ошибок P_{e1} и P_{e2} осредненных по этому ансамблю кодов и затем докажем существование в этом ансамбле отдельных кодов, удовлетворяющих таким границам для их вероятностей ошибок.

Случайный ансамбль пар кодов для такого канала, где код для передачи в направлении 1—2 из M_1 слов и в направлении 2—1 из M_2 слов, строится следующим образом. M_1 целых чисел 1, 2, ..., M_1 (сообщения для первого кода) отображается всевозможными способами в множество слов на входе X_1 длины n . Подобным же образом целые числа 1, 2, ..., M_2 (сообщение для второго кода) отображаются всевозможными способами в множество слов на входе X_2 длины n .

Если на конце 1 имеется a_1 входных букв, а на конце 2— a_2 входных букв, тогда различных входных слов длины n будет a_1^n и a_2^n .

соответственно и $a_1^{nM_1}$ отображений в случае первого кода, $a_2^{nM_2}$ в случае второго кода. Рассмотрим все пары этих кодов, общее число которых равно $a_1^{nM_1} a_2^{nM_2}$.

Каждой паре кодов придается некоторый вес, или вероятность, равная вероятности появления этой пары, если оба отображения независимы и целое число отображается в некоторое слово с вероятностью, приписанной этому слову. Таким образом, паре кодов придается вес, равный произведению вероятностей всех тех слов на входах, в которые данные числа отображаются этими кодами. Это множество пар кодов вместе с приписанными им таким образом вероятностями будем называть *случайным ансамблем пар кодов*, основанным на вероятностях $P\{X_1\}$ и $P\{X_2\}$.

Любая заданная пара кодов этого множества может быть использована для передачи информации, если уже выбран метод декодирования. Метод декодирования состоит из применения пары функций $\varphi(X_1, Y_1)$ и $\psi(X_2, Y_2)$ — частный случай функций определенных выше. Здесь X_1 пробегает всевозможные слова длины n на входе конца 1, а Y_1 — всевозможные блоки длины n , получаемых на выходе конца 1. Функция φ принимает значение от 1 до M_2 и представляет собой декодируемое сообщение для полученного Y_1 , если было передано X_1 . (Конечно, в общем случае X_1 используется в процессе декодирования, так как X_1 может влиять на Y_1 и поэтому содержит информацию, существенную для наилучшего декодирования.)

Аналогично, $\psi(X_2, Y_2)$ принимает значение от 1 до M_1 и представляет собой метод решения задачи о том, какое сообщение m_1 было передано на основе информации, имеющейся на конце 2. Заметим здесь, что декодирующие функции φ и ψ не обязательно должны быть одинаковыми для всех пар кодов из нашего ансамбля.

Отметим также, что кодирующие функции для нашего случайного ансамбля являются более специальными, чем в описанном выше общем случае. Последовательность входных букв X_1 для данного сообщения m_1 не зависит от полученных букв на конце 1. Для любого частного кода ансамбля существует точное отображение сообщений в последовательности на входе.

При заданных ансамблях пар кодов, описанных выше, и декодирующих функциях можно вычислить для каждой частной пары кодов две вероятности ошибок: P_{e1} — вероятность ошибки декодирования для первого кода и P_{e2} — для второго кода. Здесь предполагается, что различные сообщения для первого кода встречаются с вероятностями $1/M_1$; аналогичное предположение имеет место и для второго кода.

Средними вероятностями ошибок по ансамблю пар кодов назовем средние $E(P_{e1})$ и $E(P_{e2})$, при вычислении которых каждая вероятность ошибки для частного кода берется с весом, равным вероят-

ности, приписанной данной паре кодов. Требуется описать некоторый частный метод декодирования, т. е. выбор функций φ и ψ , и затем найти верхние границы для средних вероятностей ошибок по ансамблю кодов.

7. Вероятность ошибки для ансамбля кодов

Теорема 1. Пусть на основании распределений вероятностей $P\{X_1\}$ и $P\{X_2\}$ для двустороннего дискретного канала без памяти вычисляются функции распределения информации $q_{12}(Z)$ и $q_{21}(Z)$. Пусть $M_1 = \exp(R_1 n)$ и $M_2 = \exp(R_2 n)$ — произвольные целые числа, а θ_1 и θ_2 — произвольные положительные числа. Тогда случайный ансамбль пар кодов с M_1 и M_2 сообщениями соответственно имеет (при соответствующих декодирующих функциях) средние вероятности ошибок, ограниченные следующими величинами:

$$E(P_{e1}) \leq q_{12}[n(R_1 + \theta_1)] + e^{-n\theta_1}, \quad (13)$$

$$E(P_{e2}) \leq q_{21}[n(R_2 + \theta_2)] + e^{-n\theta_2}.$$

Кроме того, в ансамбле существует по крайней мере одна пара кодов, для которых индивидуальные вероятности ошибок ограничены удвоенными частями неравенств (13), т. е. удовлетворяют условиям

$$P_{e1} \leq 2q_{12}[n(R_1 + \theta_1)] + 2e^{-n\theta_1}, \quad (14)$$

$$P_{e2} \leq 2q_{21}[n(R_2 + \theta_2)] + 2e^{-n\theta_2}.$$

Теорема для двустороннего канала является обобщением теоремы 1 (из работы, цитированной на стр. 630), в которой устанавливается граница для P_e в случае одностороннего канала. Доказательство ее является обобщением доказательства для двустороннего канала.

Статистическая ситуация здесь достаточно сложна. Имеется несколько статистических событий: выбор сообщений m_1 и m_2 , выбор некоторой пары кодов из ансамбля пар кодов, и, наконец, статистика самого канала, в котором получение выходных слов Y_1 и Y_2 имеет вероятность $P\{Y_1, Y_2 | X_1, X_2\}$. Вероятности ошибок для ансамбля будем вычислять как средние по всем этим статистическим событиям.

Вначале определим системы декодирования для различных кодов ансамбля. Для данного θ_2 и для каждой пары X_1, Y_1 определим соответствующее множество слов в пространстве X_2 , обозначаемое $S(X_1, Y_1)$, следующим образом:

$$S(X_1, Y_1) = \left\{ X_2 \mid \log \frac{P\{X_1, X_2, Y_2\}}{P\{X_2\}P\{X_1, Y_1\}} > n(R_2 + \theta_2) \right\}. \quad (15)$$

Таким образом, $S(X_1, Y_1)$ представляет собой множество слов из X_2 , взаимная информация которых с фиксированной парой (X_1, Y_1) превышает некоторый уровень $n(R_2 + \theta_2)$. Аналогичным путем определим множество $S'(X_2, Y_2)$ слов из X_1 для каждой пары X_2, Y_2

$$S'(X_2, Y_2) = \left\{ X_1 \mid \log \frac{P\{X_1, X_2, Y_1\}}{P\{X_1\}P\{X_2, Y_2\}} > n(R_1 + \theta_1) \right\}. \quad (16)$$

Будем использовать множества S и S' для того, чтобы определить процесс декодирования и при отыскании верхних границ для вероятностей ошибок. Этот процесс определяется следующим образом. Для любой частной пары кодов из случайного ансамбля предположим, что поступило сообщение m_1 и оно отображается в слово X_1 на входе. Пусть Y_1 есть слово, соответствующее блоку из n букв, полученное на конце 1. Рассмотрим подмножество $S(X_1, Y_1)$ слов из X_2 . Могут встретиться следующие случаи. 1) Не имеется ни одного сообщения m_2 , отображаемого в подмножество $S(X_1, Y_1)$ для рассматриваемой пары кодов. В этом случае X_1, Y_1 декодируются, по условию, как сообщение номер один. 2) Имеется в точности одно сообщение, отображаемое в это подмножество. В этом случае оно декодируется в это сообщение. 3) Имеется более чем одно такое сообщение. В этом случае оно декодируется в то из таких сообщений, которое имеет наименьший номер.

Естественно представлять себе, что вероятность ошибок, которые требуется оценить, вычисляются следующим путем. Для каждой пары кодов подсчитываются вероятности ошибок при всевозможных сообщениях m_1 и m_2 ; осредняя их, получаем вероятности ошибок для данной пары кодов. Затем эти вероятности ошибок осредняются по всему ансамблю пар кодов с использованием приспанных каждой паре кодов весов или вероятностей. Можно, однако, изменить этот порядок осреднения. Можно рассматривать случаи, когда сообщениями являются некоторые фиксированные $\bar{m}_1$ и $\bar{m}_2$, которые отображаются в фиксированные $\bar{X}_1$ и $\bar{X}_2$, переходящие затем в слова $\bar{Y}_1$ и $\bar{Y}_2$. При этом со статистической точки зрения остается еще возможность изменения пар кодов, задаваемых теперь отображением оставшихся $M_1 - 1$ сообщений для одного кода и $M_2 - 1$ сообщений для другого. При осреднении по этому подмножеству кодов требуется показать, что вероятность того, что любое из этих оставшихся сообщений будет отображено в подмножества $S'(\bar{X}_2, \bar{Y}_2)$ и $S(\bar{X}_1, \bar{Y}_1)$, не превосходит соответственно $\exp(-n\theta_1)$ и $\exp(-n\theta_2)$.

Заметим сначала, что если X_1 принадлежит множеству $S'(\bar{X}_2, \bar{Y}_2)$, то по определению этого множества

$$\log \frac{P\{X_1, \bar{X}_2, \bar{Y}_2\}}{P\{X_1\}P\{\bar{X}_2, \bar{Y}_2\}} > n(R_1 + \theta_1), \quad (17)$$

$$P\{X_1 | \bar{X}_2, \bar{Y}_2\} > P\{X_1\} e^{n(R_1 + \theta_1)}.$$

Проссуммировав полученные неравенства по множеству значений X_1 , принадлежащих $S'(\bar{X}_2, \bar{Y}_2)$, получим

$$1 \geq \sum_{X_1 \in S'(\bar{X}_2, \bar{Y}_2)} P\{X_1 | \bar{X}_2, \bar{Y}_2\} > e^{n(R_1 + \theta_1)} \sum_{X_1 \in S'(\bar{X}_2, \bar{Y}_2)} P\{X_1\}. \quad (18)$$

Левое неравенство здесь справедливо, так как сумма вероятностей несовместных событий не может превосходить единицы. Сумму в правой части неравенства можно обозначить через $P\{S'(\bar{X}_2, \bar{Y}_2)\}$. Сопоставляя левую и правую часть нашего неравенства, получим

$$P\{S'(\bar{X}_2, \bar{Y}_2)\} < e^{-n(R_1 + \theta_1)}. \quad (19)$$

Таким образом, полная вероятность любого из множеств $S'(\bar{X}_2, \bar{Y}_1)$ ограничена выражением, содержащим n , R_1 и θ_1 , но не зависящим от фиксированных $\bar{X}_2, \bar{Y}_2$.

Теперь вспомним, что сообщения отображались во входные слова независимо, с использованием вероятностей $P\{X_1\}$ и $P\{X_2\}$. Вероятность в ансамбле пар кодов того, что некоторое фиксированное сообщение будет отображено внутрь $S'(\bar{X}_2, \bar{Y}_2)$ в точности равна $P\{S'(\bar{X}_2, \bar{Y}_2)\}$. Вероятность попасть в дополнение к этому множеству есть $1 - P\{S'(\bar{X}_2, \bar{Y}_2)\}$. Вероятность того, что *все* остальные сообщения, кроме $\bar{M}_1^1$, будут отображены в это дополнительное множество, равна

$$\begin{aligned} [1 - P\{S'(\bar{X}_2, \bar{Y}_2)\}]^{M_1 - 1} &\geq 1 - (M_1 - 1) P\{S'(\bar{X}_2, \bar{Y}_2)\} > \\ &\geq 1 - M_1 P\{S'(\bar{X}_2, \bar{Y}_2)\} \geq \\ &\geq 1 - M_1 e^{-n(R_1 + \theta_1)} = 1 - e^{-n\theta_1}. \end{aligned} \quad (20)$$

Здесь было использовано неравенство $(1 - x)^p \geq 1 - px$, соотношение (19) и, наконец, тот факт, что $M_1 = \exp(nR_1)$.

Таким образом, установлено, что в подмножестве рассмотренных случаев (тех случаев, когда сообщения $\bar{m}_1$ и $\bar{m}_2$ отображаются в $\bar{X}_1$ и $\bar{X}_2$ и затем переходят в $\bar{Y}_1$ и $\bar{Y}_2$) с вероятностью, не меньшей $1 - \exp(-n\theta_1)$, не будет других сообщений, отображаемых внутрь множества $S'(\bar{X}_2, \bar{Y}_2)$. Аналогичные вычисления показывают, что с

вероятностью, не меньшей $1 - \exp(-n\theta_2)$, не будет других сообщений, отображаемых внутрь $S(\bar{X}_1, \bar{Y}_1)$. Как уже отмечалось, эти границы не зависят от фиксированных $\bar{X}_1, \bar{Y}_1$ и $\bar{X}_2, \bar{Y}_2$.

Найдем теперь границу для вероятности того, что поступившее сообщение $\bar{m}_1$ будет отображено внутрь подмножества $S'(\bar{X}_2, \bar{Y}_2)$. Напомним, что из определения $q_{12}(Z)$ следует,

$$q_{12}[n(R_1 + \theta_1)] = P \left\{ \log \frac{P\{X_1, X_2, Y_2\}}{P\{X_1\}P\{X_2, Y_2\}} \leq n(R_1 + \theta_1) \right\}. \quad (21)$$

В ансамбле пар кодов некоторое сообщение, скажем $\bar{m}_1$, отображается в слова X_1 с вероятностями, в точности равными $P\{X_1\}$. Следовательно, вероятность того, что поступившее сообщение было отображено внутрь $S'(\bar{X}_2, \bar{Y}_2)$, зависящая от ансамбля пар кодов сообщения и статистики канала, в точности равна $1 - q_{12}[n(R_1 + \theta_1)]$.

Вероятность того, что поступившее сообщение было отображено *вне* $S'(\bar{X}_2, \bar{Y}_2)$, равна поэтому $q_{12}[n(R_1 + \theta_1)]$, а вероятность отображения любого другого сообщения *внутри* $S'(\bar{X}_2, \bar{Y}_2)$ ограничена, как показано выше, величиной $\exp(-n\theta_1)$. Вероятность того, что истинно *хотя бы одно* из этих событий, ограничена поэтому выражением $q_{12}[n(R_1 + \theta_1)] + \exp(-n\theta_1)$; но тогда это же выражение является границей и для $E(P_{e1})$, так как если ни одно из описанных событий не наступило, то декодирование с применением нашего метода даст правильный результат.

Эти же рассуждения с переменной индексов дают соответствующую границу для $E(P_{e2})$. Этим заканчивается доказательство первой части теоремы.

Приступая к доказательству последнего утверждения теоремы, вначале докажем простую комбинаторную лемму, которая будет полезна не только здесь, но и в других вопросах теории кодирования.

Лемма. *Предположим, что имеется некоторое множество объектов $B_1, B_2, \dots, B_n$ с приписанными им вероятностями $P_1, P_2, \dots, P_n$ и некоторое количество числовых свойств (функций) этих объектов $f_1, f_2, \dots, f_d$. Все они являются неотрицательными, $f_i(B_j) \geq 0$, и нам известны средние A_i этих свойств объектов*

$$\sum_j P_j f_i(B_j) = A_i, \quad i = 1, 2, \dots, d. \quad (22)$$

Тогда существует объект B_p , для которого

$$f_i(B_p) \leq dA_i, \quad i = 1, 2, \dots, d. \quad (23)$$

Вообще, для любого множества чисел $K_i > 0$, таких, что $\sum_{i=1}^d (1/K_i) \leq 1$, существует объект B_p , для которого

$$f_i(B_p) \leq K_i A_i, \quad i = 1, 2, \dots, d. \quad (24)$$

Доказательство. Из второй части леммы вытекает первая, если положить $K_i = d$. Чтобы доказать вторую часть, обозначим через Q_i суммарную вероятность объектов B , для которых $f_i(B) > K_i A_i$. Теперь среднее $A_i > Q_i K_i A_i$, так как $Q_i K_i A_i$ есть вклад в общую сумму тех B_i , для которых $f_i(B) > K_i A_i$, а для всех оставшихся B значения $f_i \geq 0$. Таким образом,

$$Q_i < \frac{1}{K_i}, \quad i = 1, 2, \dots, d. \quad (25)$$

Общая вероятность Q объектов, для которых нарушается какое-либо из наших условий, меньше или равна сумме всех частных Q_i , так что

$$Q < \sum_{i=1}^d \frac{1}{K_i} \leq 1. \quad (26)$$

Таким образом, имеется по крайней мере один объект, для которого не нарушается ни одно из наших условий, чем и заканчивается доказательство.

Предположим, например, нам известно, что в комнате находится некоторое число людей, средний возраст которых составляет 40 лет, а средний рост 5 футов. Здесь $d=2$ и, используя более простую формулировку теоремы, можно утверждать, что в комнате имеется некто не старше 80 лет и не выше 10 футов, даже если в комнате находятся пожилые карлики и молодые баскетболисты. Полагая $K_1 = 8/3$, $K_2 = 8/5$, можно утверждать, что присутствует человек не выше 8 футов и не старше $106^{2/3}$ лет.

Теперь, возвращаясь к доказательству теоремы 1, можно установить последнее предположение. Имеем некоторое множество объектов — пар кодов, и два свойства каждого объекта — вероятность ошибки P_{e1} для кода в направлении 1—2 и вероятность ошибки P_{e2} для кода в направлении 2—1. Они являются неотрицательными и их средние ограничены выражениями, содержащимися в первой части теоремы 1. Из комбинаторной леммы следует, что существует по крайней мере одна фиксированная пара кодов, для которой одновременно

$$\begin{aligned} P_{e1} &\leq 2 \{Q_{12} [n(R_1 + \theta_1)] + e^{-n\theta_1}\}, \\ P_{e2} &\leq 2 \{Q_{21} [n(R_2 + \theta_2)] + e^{-n\theta_2}\}. \end{aligned} \quad (27)$$

Этим заканчивается доказательство теоремы 1.

Легко видеть, что эта теорема доказывает существование пар кодов со скоростями передачи R_1 и R_2 , произвольно близкими к средней взаимной информации для букв R_{12} и R_{21} при любых заданных $P\{x_1\}$ и $P\{x_2\}$ и с произвольно малыми вероятностями ошибок. Действительно, пусть $R_{12} - R_1 = R_{21} - R_2 = \varepsilon > 0$ и в формулировке теоремы $\theta_1 = \theta_2 = \varepsilon/2$. Так как $Q_{12}[n(R_{12} - \varepsilon/2)] \rightarrow 0$ и притом экспоненциально быстро относительно n (как функция распределения суммы соответствующим образом нормированных независимых случайных величин¹⁾), граница для P_{e1} стремится к нулю экспоненциально по n . Аналогичные рассуждения можно провести и для границы P_{e2} . Выбирая затем возрастающие по n последовательности из M_1 и M_2 , при передаче которых аппроксимируются снизу желаемые скорости R_1 и R_2 , получим нужный результат, который можно сформулировать следующим образом.

Теорема 2. Пусть дан двусторонний канал 'без памяти с вероятностями входных букв $P\{x_1\}$ и $P\{x_2\}$ и средними взаимными информацией по обоим направлениям

$$R_{12} = E \left(\log \frac{P\{x_1|x_2, y_2\}}{P\{x_1\}} \right), \quad R_{21} = E \left(\log \frac{P\{x_2|x_1, y_1\}}{P\{x_2\}} \right). \quad (28)$$

Тогда для данного $\varepsilon > 0$ существует некоторая пара кодов для всех достаточно больших длин блоков n со скоростями передачи по обоим направлениям большими, чем $R_{12} - \varepsilon$ и $R_{21} - \varepsilon$, соответственно и вероятностями ошибок

$$P_{e1} \leq \exp[-A(\varepsilon)n], \quad P_{e2} \leq \exp[-A(\varepsilon)n],$$

где $A(\varepsilon)$ — есть некоторое положительное число, не зависящее от n .

Придавая входным буквам различные вероятности и используя эту теорему, можно получить различные точки области пропускной способности. Конечно, чтобы получить на основе этой теоремы наилучшие возможные скорости, надо искать максимум этих скоростей. Это на самом деле можно сделать, используя метод Лагранжа максимизации $R_{12} + \lambda R_{21}$ при различных положительных λ .

8. Выпуклая оболочка G_1 как внутренняя граница области пропускной способности

В дополнение к скоростям, полученным вышеизложенным методом, можно построить коды, которые являются *смесями* кодов, полученных с помощью описанного процесса. Предположим, что

¹⁾ Крамер Г.. Одна предельная теорема теории вероятностей, *Успехи мат. наук* (ст. серия), 10 1944, 166. — Прим. ред.

при одном выборе распределений вероятностей $P\{x_1\}$, $P\{x_2\}$ получаются средние взаимные информации R_{12} , R_{21} , а при втором выборе $P'\{x_1\}$, $P'\{x_2\}$ — информации R'_{12} , R'_{21} . Тогда можно найти, основываясь на первом выборе распределений, код (достаточно большой) длины n с вероятностями ошибок больше δ и со скоростями, отличающимися от R_{12} , R_{21} не более чем на ϵ ; и, основываясь на распределениях $P'\{x_1\}$, $P'\{x_2\}$, второй код длины n' , с теми же δ и ϵ . Рассмотрим теперь код длины $n+n'$ с $M_1M'_1$ словами в прямом направлении и $M_2M'_2$ — в обратном, состоящий из всевозможных слов первого кода, за которыми следуют всевозможные слова второго кода для того же направления. Скорости передачи для этого кода R_1^* и R_2^* равны взвешенным средним скоростей передачи для первоначальных кодов $[R_1^* = nR_1/(n+n') + n'R'_1/(n+n')$; $R_2^* = nR_2/(n+n') + n'R'_2/(n+n')]$ и, следовательно, отличаются не более чем на ϵ от взвешенных средних R_{12} и R'_{12} $|R_1^* - nR_{12}/(n+n') - n'R'_{12}/(n+n')| < \epsilon$; аналогичное неравенство справедливо и для R_2^* .

Кроме того, вероятность ошибки при этом коде не превосходит 2δ , так как вероятность наступления какого-либо из двух событий (ошибки в какой-либо из двух частей кода) не превосходит суммы вероятностей этих событий. Можно построить такой смешанный код для любых достаточно больших n и n' . Следовательно, выбирая их достаточно большими, можно аппроксимировать любое среднее взвешенное от данных скоростей так, что при этом вероятность ошибки будет экспоненциально быстро стремиться к нулю. Отсюда следует, что можно присоединить к множеству точек, полученных с помощью задания различных распределений вероятностей на входных буквах, все точки, принадлежащие выпуклой оболочке этого множества. Этот метод действительно добавляет в некоторых случаях новые точки, как, например, в канал с несовместимой передачей по двум направлениям (табл. 1). Смешивая в равной пропорции коды для распределений вероятностей на входных буквах, дающих точки (0,1) и (1,0), получаем точку (1/2, 1/2). Пара скоростей, соответствующая этой смеси, не может быть получена ни при каком распределении вероятностей для одной буквы. Изложенное можно суммировать в виде следующей теоремы:

Теорема 3. Пусть G_1 есть выпуклая оболочка множества точек (R_{12}, R_{21})

$$R_{12} = E \left(\log \frac{P\{x_1 | x_2, y_2\}}{P\{x_1\}} \right), \quad R_{21} = E \left(\log \frac{P\{x_2 | x_1, y_1\}}{P\{x_2\}} \right), \quad (29)$$

где $P\{x_1\}$ и $P\{x_2\}$ пробегают различные значения распределения вероятностей. Все точки G_1 принадлежат области пропускной способности канала. Для любой точки (R_1, R_2) в G_1 и любого $\epsilon > 0$ можно найти коды, скорости передачи для которых отличаются

соответственно от R_1 и R_2 не более чем на ϵ , а вероятности ошибок по обоим направлениям не превосходят $\exp[-A(\epsilon)n]$ для всех достаточно больших n и некоторого положительного $A(\epsilon)$.

Отметим, что выпуклая оболочка G_1 , фигурирующая в данной теореме, является замкнутым множеством (содержащим все свои предельные точки). Это вытекает из непрерывности R_{12} и R_{21} как функций от распределений вероятностей $P\{x_1\}$ и $P\{x_2\}$. Кроме того, если G_1 содержит некоторую точку (R_1, R_2) , то она содержит также и ее проекции $(R_1, 0)$ и $(0, R_2)$. Это сейчас будет доказано.

Высказанное утверждение станет очевидным, если удастся показать, что проекции любой точки, полученной исходя из некоторого распределения вероятностей на входных буквах, также принадлежит G_1 . Чтобы показать это, предположим, что распределения $P\{x_1\}$ и $P\{x_2\}$ дают точку (R_{12}, R_{21}) . Тогда R_{12} является средним для различных частных значений R_{12} , получающихся, когда x_2 принимает различные фиксированные значения. Таким образом,

$$R_{12} = \sum_{x_2} P\{x_2\} \sum_{x_1, y_2} P\{x_1, y_2 | x_2\} \log \frac{P\{x_1 | x_2, y_2\}}{P\{x_1\}}. \quad (30)$$

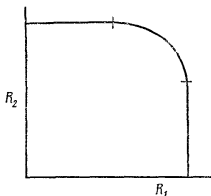
Отсюда следует, что существует некоторое фиксированное значение x_2 , скажем x_2^* , для которого внутренняя сумма по крайней мере не меньше среднего, т. е. для которого

$$\begin{aligned} \sum_{x_1, y_2} P\{x_1, y_2 | x_2^*\} \log \frac{P\{x_1 | x_2^*, y_2\}}{P\{x_1\}} &> \\ &\geq \sum_{x_2} P\{x_2\} \sum_{x_1, y_2} P\{x_1, y_2 | x_2\} \log \frac{P\{x_1 | x_2, y_2\}}{P\{x_1\}}. \end{aligned} \quad (31)$$

Приписывая вероятности $P\{x_1 | x_2^*\}$ буквам x_1 и положив $P\{x_2\} = 1$, если $x_2 = x_2^*$ и 0 в остальных случаях, получим некоторую точку на горизонтальной оси, совпадающую с проекцией данной точки (R_{12}, R_{21}) или лежащую правее ее. Поступая аналогично, можно найти точку x_1^* , такую, что, приписывая буквам x_2 вероятности $P\{x_2 | x_1^*\}$ и положив $P\{x_1\} = 1$, получим некоторую точку на вертикальной оси, совпадающую с проекцией (R_{12}, R_{21}) или лежащую выше ее. Заметим также, что, положив $P\{x_1^*\} = 1$, $P\{x_2^*\} = 1$, получим точку $(0, 0)$. Следовательно, смешивая соответствующим образом коды, полученные для этих четырех распределений вероятностей, можно аппроксимировать любую точку четырехугольника, определяемого соответствующими парами скоростей и, в частности, любую точку прямоугольника с вершиной в точке (R_{12}, R_{21}) . Из этих замечаний следует, что выпуклая оболочка G_1 — область, типичная форма которой показана на рис. 7. Она ограничена горизонтальным отрезком, выпуклой кривой, вертикальным отрезком

и двумя отрезками осей; любая из этих частей может равняться нулю. Выпуклая оболочка G_1 , как было показано, лежит внутри области пропускной способности. Назовем ее *внутренней границей*.

Довольно интересно попытаться точнее оценить скорость уменьшения вероятности ошибки с возрастанием длины кода n . Этот



Р и с. 7.

вопрос рассматривается в приложении и приводит к обобщению теоремы 2 из работы автора¹). Оценку мы получим, используя логарифм производящей функции для моментов.

9. Внешняя граница для области пропускной способности

Хотя в некоторых случаях выпуклая оболочка G_1 — внутренняя граница, определенная выше, — действительно является областью пропускной способности, однако это бывает далеко не всегда. Галлагер (R. G. Gallager) с помощью тонких вычислений показал, что для двоичного умножающего канала внутренняя граница находится строго внутри области пропускной способности. Однако при *частичном* обращении теоремы 3 можно получить некоторую внешнюю границу для области пропускной способности. Пусть имеется некоторый код, начинающийся в нулевой момент времени с сообщений m_1 и m_2 на обоих концах. Пусть Y_1 и Y_2 — полученные блоки на обоих концах канала (т. е. последовательности из n букв) после n операций в канале и пусть x_1, x_2, y_1 и y_2 — следующие переданные и полученные буквы. Рассмотрим изменение в неопределенности сообщений на обоих концах канала, даваемое следующей полученной буквой. Например, на конце 2 это изменение есть

¹) Certain results in coding theory for noisy channels, *Information and control*, 1, 1957, 6. (Русский перевод см. стр. 509 данного сборника.—Прим. ред.)

(некоторые очевидные преобразования опущены)

$$\begin{aligned}\Delta &= H(m_1 | m_2, Y_2) - H(m_1 | m_2, Y_2, y_2) = \\ &= E \left[\log \frac{P\{m_2, Y_2\}}{P\{m_1, m_2, Y_2\}} \right] - E \left[\log \frac{P\{m_2, Y_2, y_2\}}{P\{m_1, m_2, Y_2, y_2\}} \right] = \\ &= E \left[\log \frac{P\{y_2 | m_1, m_2, Y_2\} P\{y_2 | x_2\}}{P\{y_2 | x_2\} P\{y_2 | Y_2, m_2\}} \right].\end{aligned}\quad (32)$$

Имеет место соотношение $H(y_2 | m_1, m_2, Y_2) \geq H(y_2 | m_1, m_2, Y_1, Y_2) = H(y_2 | x_1, x_2)$, так как при добавлении введенных условий энтропия не может возрасти и $P\{y_2 | m_1, m_2, Y_1, Y_2\} = P\{y_2 | x_1, x_2\}$.

Аналогично $H(y_2 | x_2) \geq H(y_2 | Y_2, m_2)$, так как x_2 есть некоторая функция от Y_2 и m_2 , получаемая при кодировании. Отсюда

$$\Delta \leq E \left(\log \frac{P\{y_2 | x_1, x_2\}}{P\{y_2 | x_2\}} \right) + H(y_2 | Y_2, m_2) - H(y_2 | x_2), \quad (33)$$

$$\begin{aligned}\Delta &\leq E \left(\log \frac{P\{y_2 | x_1, x_2\}}{P\{y_2 | x_2\}} \right) = E \left(\log \frac{P\{y_2, x_1, x_2\} P\{x_2\}}{P\{x_2, y_2\} P\{x_1, x_2\}} \right) = \\ &= E \left(\log \frac{P\{x_1 | x_2, y_2\}}{P\{x_1 | x_2\}} \right).\end{aligned}\quad (34)$$

Полученный результат действительно приводит к обращению теоремы 1, если только случайные величины x_1 и x_2 независимы, так как в этом случае последнее выражение сводится к виду $E[\log(P\{x_1 | x_2, y_2\}/P\{x_1\})]$. К сожалению, в общем случае кодирования x_1 и x_2 не являются обязательно независимыми. Действительно, следующие x_1 и x_2 могут быть функционально связаны с полученными X и Y и, следовательно, быть зависимыми.

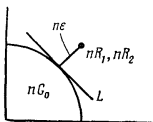
Однако можно получить по крайней мере внешнюю границу для области пропускной способности канала. А именно приведенное выше неравенство вместе с аналогичным неравенством для второго конца канала показывает, что изменение вектора неопределенности, даваемое другими полученными буквами, должно быть вектором, компоненты которого не превосходят величин

$$E \left(\log \frac{P\{x_1 | x_2, y_2\}}{P\{x_1 | x_2\}} \right), \quad E \left(\log \frac{P\{x_2 | x_1, y_1\}}{P\{x_2 | x_1\}} \right) \quad (35)$$

для некоторого $P\{x_1, x_2\}$. Таким образом, это изменение вектора содержится в выпуклой оболочке G_0^2 всех таких векторов (когда изменяется $P\{x_1, x_2\}$).

Для кода длины n общее изменение неопределенности от начала и до конца блока не может превосходить суммы n векторов из нашей выпуклой оболочки. Таким образом, эта сумма будет находиться внутри выпуклой оболочки nG_0 , т. е. внутри G_0 , расширенной в n раз.

Предположим теперь, что нам дан код, имеющий скорости передачи $R_1 = (1/n) \log M_1$ и $R_2 = (1/n) \log M_2$. Тогда первоначальная неопределенность сообщений есть nR_1 и nR_2 . Предположим далее, что точка (nR_1, nR_2) лежит вне выпуклой оболочки nG_0 , причем кратчайшее расстояние от этой точки до оболочки равно $n\epsilon$ (см. рис. 8). Построим прямую L , проходящую через ближайшую к (nR_1, nR_2) точку области и nG_0 перпендикулярно к кратчайшему отрезку, соединяющему nR_1, nR_2 и nG_0 ; тогда область



Р и с. 8.

nG_0 будет находиться с одной стороны от прямой L (здесь используется то обстоятельство, что nG_0 есть выпуклая область). Отсюда ясно, что для любой точки $(nR_1^* \text{ и } nR_2^*)$, лежащей с той же стороны от линии L , что и область nG_0 и, в частности, для любой точки самой области nG_0 , имеем $|nR_1 - nR_1^*| + |nR_2 - nR_2^*| \geq n\epsilon$ (так как наименьшее расстояние равно $n\epsilon$). Кроме того, по крайней мере одна из величин $nR_1 - nR_1^*$ и $nR_2 - nR_2^*$ будет не меньше $n\epsilon/\sqrt{2}$. (В прямоугольном треугольнике по крайней мере один катет не меньше гипотенузы, деленной на $\sqrt{2}$.)

Таким образом, если после n -кратного использования канала пара скоростей передачи (R_1, R_2) находилась вне выпуклой оболочки G_0 на расстоянии ϵ , то по крайней мере одна из окончательных неопределенностей будет не меньше $\epsilon/\sqrt{2}$, где все неопределенности измеряются за секунду. Таким образом, неопределенности за секунду для скоростей передачи, находящихся вне G_0 на расстоянии ϵ , ограничены снизу независимо от длины кода n . Отсюда вытекает, что вероятность ошибки также ограничена снизу, т. е. по крайней мере один из пары кодов имеет вероятность ошибки, не меньшую $f(\epsilon) > 0$ и не зависящую от n , как это показано в приложении к работе автора¹⁾.

Суммируя вышеизложенное, можно сказать, что область пропускной способности G содержится внутри выпуклой оболочки G_0

¹⁾ Shannon C., Channels with side information at the transmitter. (Русский перевод см. стр. 497 данного сборника. — Прим. ред.)

точек (R_{12}, R_{21}) , где

$$\begin{aligned} R_{12} &= E \left[\log \frac{P\{x_1 | x_2, y_2\}}{P\{x_1 | x_2\}} \right], \\ R_{21} &= E \left[\log \frac{P\{x_2 | x_1, y_1\}}{P\{x_2 | x_1\}} \right] \end{aligned} \quad (36)$$

и совместные распределения вероятностей $P\{x_1, x_2\}$ произвольны.

Таким образом, внутренняя граница G_1 и внешняя граница G_0 находятся с помощью одного и того же процесса: задаем распределения вероятностей на входных буквах, вычисляем получающиеся средние взаимные информации R_{12} и R_{21} и берем выпуклую оболочку полученных точек. Различие состоит только в том, что при вычислении внешней границы берутся всевозможные совместные распределения вероятностей $P\{x_1, x_2\}$, в то время как при вычислении внутренней границы рассмотрение ограничивается лишь случаем независимости $P\{x_1\}P\{x_2\}$.

Выясним теперь некоторые свойства внешней границы.

10. Выпуклость вниз R_{12} и R_{21} как функций от $P\{x_1, x_2\}$

Теорема 4. *Вычисляемые на основе переходных вероятностей $P\{y_1, y_2 | x_1, x_2\}$ канала K скорости*

$$\begin{aligned} R_{12} &= E \left[\log \frac{P\{y_2 | x_1, x_2\}}{P\{y_2 | x_2\}} \right], \\ R_{21} &= E \left[\log \frac{P\{y_1 | x_1, x_2\}}{P\{y_1 | x_1\}} \right] \end{aligned} \quad (37)$$

являются выпуклыми функциями от совместного распределения вероятностей $P\{x_1, x_2\}$ на входах канала. Например,

$$\begin{aligned} R_{12}(P_1\{x_1, x_2\}/2 + P_2\{x_1, x_2\}/2) &\geq \\ &\geq R_{12}(P_1\{x_1, x_2\})/2 + R_{12}(P_2\{x_1, x_2\})/2. \end{aligned}$$

Это свойство является обобщением подобного же свойства для односторонних каналов, рассматриваемого в работе автора¹⁾. Для того чтобы доказать теорему, достаточно в соответствии с известными свойствами выпуклых функций показать, что

$$\begin{aligned} R_{12} \left(\frac{1}{2} P_1\{x_1, x_2\} + \frac{1}{2} P_2\{x_1, x_2\} \right) &\geq \\ &\geq \frac{1}{2} R_{12}(P_1\{x_1, x_2\}) + \frac{1}{2} R_{12}(P_2\{x_1, x_2\}). \end{aligned} \quad (38)$$

¹⁾ Shannon C., Geometrische Deutung einiger Ergebnisse bei der Berechnung der Kanalkapazität. (Русский перевод см. стр. 488 данного сборника.—Прим. ред.)

Но $R_{12}(P_1\{x_1, x_2\})$ и $R_{12}(P_2\{x_1, x_2\})$ могут быть записаны так:

$$R_{12}(P_1\{x_1, x_2\}) = \sum_{x_2} P_1\{x_2\} \sum_{x_1, y_2} P_1\{x_1, y_2 | x_2\} \log \frac{P_2\{y_2 | x_1, x_2\}}{P_1\{y_2 | x_2\}}, \quad (39)$$

$$R_{12}(P_2\{x_1, x_2\}) = \sum_{x_2} P_2\{x_2\} \sum_{x_1, y_2} P_2\{x_1, y_2 | x_2\} \log \frac{P_2\{y_2 | x_1, x_2\}}{P_2\{y_2 | x_2\}}. \quad (40)$$

Здесь подстрочный индекс 1 у вероятностей соответствует тем распределениям, которые порождаются совместными распределениями $P_1\{x_1, x_2\}$ на входах канала, аналогичное справедливо и для индекса 2. Внутреннюю сумму

$$\sum_{x_1, y_2} P_1\{x_1, y_2 | x_2\} \log (P_1\{y_2 | x_1, x_2\} / P_1\{y_2 | x_2\})$$

можно толковать как скорость передачи по каналу от x_1 к y_2 при условии, что фиксировано некоторое значение x_2 и что распределение вероятностей на буквах x_1 является условным, получающимся из совместного распределения $P_1\{x_1, x_2\}$.

Соответствующая внутренняя сумма с распределением вероятностей $P_2\{x_1, x_2\}$, имеющим вид

$$\sum_{x_1, y_2} P_2\{x_1, y_2 | x_2\} \log (P_2\{y_2 | x_1, x_2\} / P_2\{y_2 | x_2\}),$$

может рассматриваться как условная скорость передачи букв x_2 для такого же одностороннего канала, но с распределением вероятностей $P_2\{x_1 | x_2\}$ на входных буквах.

Идя по этому пути далее, можно применить результаты работы¹⁾, относящиеся к свойству выпуклости. В частности, взвешенное среднее от этих скоростей, где весами служат распределения $P_1\{x_2\} / (P_1\{x_2\} + P_2\{x_2\})$ и $P_2\{x_2\} / (P_1\{x_2\} + P_2\{x_2\})$, превосходит скорость передачи в соответствующем одностороннем канале, распределение вероятностей на входе которого является таким же средним взвешенным от обоих данных распределений. Это среднее взвешенное равно

$$\begin{aligned} P_3\{x_1, x_2\} &= \frac{P_1\{x_2\}}{P_1\{x_2\} + P_2\{x_2\}} P_1\{x_1 | x_2\} + \frac{P_2\{x_2\}}{P_1\{x_2\} + P_2\{x_2\}} P_2\{x_1 | x_2\} = \\ &= \frac{1}{2} \frac{1}{P_1\{x_2\} + P_2\{x_2\}} \cdot 2(P_1\{x_1, x_2\} + P_2\{x_1, x_2\}). \end{aligned} \quad (41)$$

Таким образом, сумма двух соответствующих слагаемых из неравенства (38) (с одним и тем же x_2) мажорируется величиной $P_1\{x_2\} + P_2\{x_2\}$, умноженной на скорость передачи по соответствующему одностороннему каналу, распределение на входе кото-

¹⁾ См. работу, указанную в примечании на стр. 643. — *Прим. ред.*

рого задают выписанные выше осредненные вероятности. Эта последняя скорость при подстановке осредненных вероятностей превращается в

$$\sum_{x_1, y_2} P_3 \{x_1, y_2 | x_2\} \log \frac{P_1 \{y_2 | x_1, x_2\}}{P_3 \{y_2 | x_2\}}, \quad (42)$$

где подстрочный индекс 3 соответствует вероятностям, полученным при помощи $P_3 \{x_1, x_2\} = (P_1 \{x_1, x_2\} + P_2 \{x_1, x_2\})/2$. Другими словами, суммы (39) и (40) (включая и первое суммирование по x_2) мажорируются выражением

$$\begin{aligned} \sum_{x_2} (P_1 \{x_2\} + P_2 \{x_2\}) \sum_{x_1, y_2} P_3 \{x_1, y_2 | x_2\} \log \frac{P_3 \{y_2 | x_1, x_2\}}{P_3 \{y_2 | x_2\}} = \\ = 2 \sum_{x_1, x_2, y_2} P_3 \{x_1, y_2, x_2\} \log \frac{P_3 \{y_2 | x_1, x_2\}}{P_3 \{y_2 | x_2\}}. \end{aligned} \quad (43)$$

Это и есть нужное нам утверждение теоремы.

11. Приложения свойства выпуклости Каналы с симметричной структурой

Теорема 4 бывает полезна при явных оценках внешней границы для конкретных каналов. Во-первых, заметим, что $R_{12} + \lambda R_{21}$ как функция от $P\{x_1, x_2\}$ при положительном λ также является выпуклой вниз функцией. Следовательно, любой локальный максимум является абсолютным и числовое исследование при отыскании этого максимума с помощью метода множителей Лагранжа тем самым упрощается.

Кроме того, вывод о выпуклости оказывается очень полезным при отыскании такого максимума для каналов, обладающих некоторой «симметрией». Пусть, например, дан канал с множеством переходных вероятностей $P\{y_1, y_2 | x_1, x_2\}$, обладающих следующим свойством: существует такая перестановка (перенумерация) входных букв x_1 и выходных y_1 и y_2 , при которой, скажем, переставляются две первые буквы в x_1 , но совокупность вероятностей $P\{y_1, y_2 | x_1, x_2\}$ остается той же самой. Если теперь некоторое фиксированное распределение $P\{x_1, x_2\}$ дает скорости R_{12} и R_{21} из внешней границы, тогда, применяя ту же самую перестановку алфавита x для распределения $P\{x_1, x_2\}$, получим новое распределение вероятностей, которое, однако, будет давать те же самые скорости R_{12} и R_{21} из внешней границы. Если взять среднее этих двух распределений вероятностей, то на основании свойства выпуклости получим новое распределение, которое дает не меньшие значения R_{12} и R_{21} . В этом усредненном распределении при любом фиксированном x_2 первые две буквы из алфавита x_1 имеют равные вероятности.

Другими словами, в этом случае некоторое распределение, скажем $P\{x_1, x_2\}$, максимизирующее $R_{12} + \lambda R_{21}$ и представленное в виде матрицы, должно иметь две одинаковые первые строки.

Если канал обладает достаточно симметричной структурой, так что при перестановке любой пары букв в алфавите x_1 можно так переставить буквы в алфавитах x_2 , y_1 и y_2 , что $P\{y_1, y_2|x_1, x_2\}$ сохраняется, то будет существовать некоторое максимизирующее распределение $P\{x_1, x_2\}$, в котором все строки одинаковы. В этом случае оно является функцией только от x_2 : $P\{x_1, x_2\} = P\{x_2\}/\alpha$, где α есть число букв в алфавите x_1 . Таким образом, максимум по всем распределениям $P\{x_1, x_2\}$ в действительности достигается при независимых x_1 и x_2 . Другими словами, в случае, когда имеет место симметрия относительно любых перестановок в алфавите x_1 , внутренняя и внешняя границы области пропускной способности совпадают. Это позволяет выделить важный класс каналов, для которых область пропускной способности может быть определена сравнительно легко.

Примером этого типа каналов служит канал с такими переходными вероятностями, что все входные и выходные буквы являются двоичными $y_1 = x_2$ (так что имеется двоичный канал без шума, осуществляющий передачу от конца 2 к концу 1). Если $x_2 = 0$, тогда $y_2 = x_1$, если же $x_2 = 1$, то y_2 с вероятностями $1/2$ принимает значения 0 или 1. Другими словами, если $x_2 = 0$, то двоичный канал для передачи в прямом направлении является каналом без шума, если же $x_2 = 1$, то на приемном конце получается только шум. Заметим здесь, что если сделать перестановку букв алфавита x_1 и одновременно переставить буквы алфавита y_2 , то канал останется прежним и условные вероятности не изменятся. Из проведенного выше анализа видно, что внутренняя и внешняя границы совпадают и дают область пропускной способности канала. Кроме того, граничные точки будут здесь получаться при распределениях $P\{x_1, x_2\}$, матрица которых имеет равные строки, как показано в табл. II.

Таблица II

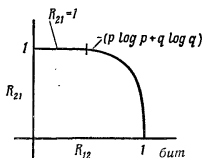
		x_2	
		0	1
x_1	0	$p/2$	$q/2$
	1	$p/2$	$q/2$

Для фиксированного p такое распределение дает скорости

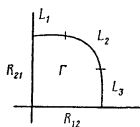
$$R_{12} = p; \quad R_{21} = -(p \log p + q \log q). \quad (44)$$

Эти формулы можно получить непосредственной подстановкой в формулы для R_{12} и R_{21} или же, заметив, что при передаче по направлению 1—2 канал действует подобно стирающему каналу, а по направлению 2—1 ведет себя как двоичный канал без шума при неравных вероятностях входных букв. Это дает область пропускной способности, изображенную на рис. 9.

Эти искусственные приемы, основанные на использовании перестановок и свойств симметрии, прилагаются к вычислению области пропускной способности во многих различных вариантах. Например,



Р и с. 9.



Р и с. 10.

если при любых перестановках букв в обоих алфавитах сохраняются переходные вероятности, то максимизирующее распределение должно иметь одинаковыми и строки и столбцы, и входы идентичны. Тогда $P\{x_1, x_2\} = 1/\alpha c$, где α и c есть соответственно число букв в алфавитах x_1 и x_2 . В этом случае всевозможные достижимые точки (R_{12}, R_{21}) мажорируются фиксированной точкой, получаемой из нашего равномерного распределения вероятностей. Другими словами, в случае симметрии канала относительно всевозможных перестановок букв в алфавитах x_1 и x_2 область ее пропускной способности является прямоугольником.

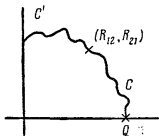
Примером этого типа служит канал, изображенный на рис. 2, определяемый соотношениями $y_1 = y_2 = x_1 \oplus x_2$, где $\oplus$ означает сложение по модулю 2.

12. Характер области, выделяемой внешней границей

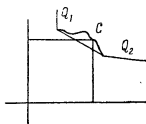
Теперь используем свойство выпуклости, чтобы установить некоторые факты относительно множества Γ точек (R_{12}, R_{21}) , которые могут быть получены при всевозможных распределениях $P\{x_1, x_2\}$ для данного канала K и выпуклая оболочка которых есть G_0 . Докажем, что множество Γ в действительности уже само выпукло и поэтому совпадает с G_0 и что оно состоит из всех внутренних и граничных точек области, типичная форма которой изображена на рис. 10; она ограничена горизонтальным отрезком L_1 , выпуклым отрезком

кривой L_2 , вертикальным отрезком L_3 и двумя отрезками координатных осей. Таким образом, G_0 имеет структуру, аналогичную G_I .

Пусть некоторое распределение $P\{x_1, |x_2\}$ дает точку (R_{12}, R_{21}) . Здесь R_{12} , как было показано ранее, есть среднее всевозможных скоростей R_{12} , которые могут быть получены при различных фиксированных значениях x_2 , используемых с вероятностью 1, и при задании условного распределения $P\{x_1 | x_2\}$ на буквах x_1 . При осреднении весами служат множители $P\{x_2\}$. Отсюда следует, что при некотором фиксированном x_2 можно получить величину R_{12} , по крайней



Р и с. 11.



Р и с. 12.

мере не меньшую, чем наше среднее взвешенное. Если это фиксированное значение x_2 обозначить x_2^* , то распределение $P\{x_1 | x_2^*\}$ дает значение R_{12} , по крайней мере не меньшее первоначального, и в то же время дает значение $R_{21} = 0$. Этим методом можно получить на рис. 11 точку, принадлежащую области Γ , которая совпадает с проекцией данной точки (R_{12}, R_{21}) или лежит правее ее, как это изображено на рисунке (точка Q).

Рассмотрим теперь смеси указанных выше двух распределений вероятностей, т. е. распределение вида $\lambda P\{x_1, x_2\} + (1-\lambda)P\{x_1 | x_2^*\}$. Здесь λ непрерывно изменяется от 0 до 1. Так как R_{12} и R_{21} являются непрерывными функциями от распределения вероятностей, то таким образом получаем непрерывную кривую C , идущую от данной точки (R_{12}, R_{21}) к точке Q . Кроме того, эта кривая целиком лежит правее и выше соответствующего отрезка прямой, соединяющего эти две точки. Это следует из свойства выпуклости выражений R_{12} и R_{21} . Аналогичным путем построим изображенную на рисунке кривую C' , состоящую из точек, принадлежащих Γ и находящихся на горизонтальном отрезке прямой, проходящем через данную точку (R_{12}, R_{21}) , или выше его.

Для каждой из точек кривых C и C' рассмотрим смеси соответствующих распределений вероятностей с распределением $P\{x_1^*, x_2^*\} = 1$ (вероятность всех остальных пар равна 0). Это последнее распределение дает точку $(0,0)$. Будем постепенно менять от 0 до 1 вес при распределении, сосредоточенном в $(0,0)$. Тогда кривая, состоящая из полученных при осреднении точек, будет непрерывно изменяться, начиная от точки на кривой CC' и кончая точкой $(0,0)$.

Концевые точки получающихся кривых при этой операции остаются на отрезках координатных осей. Следовательно, опираясь на известные топологические результаты, можно утверждать, что получающиеся кривые полностью заполняют область, ограниченную C , C' и координатными осями и, в частности, *покрывают прямоугольник со сторонами, параллельными осям, и с вершиной в исходной точке (R_{12}, R_{21}) .*

Покажем теперь, что множество точек Γ есть выпуклое множество. Предположим, что изображенные на рис. 12 точки Q_1 и Q_2 есть две точки, которые получаются при распределениях $P_1\{x_1, x_2\}$ и $P_2\{x_1, x_2\}$.

Составляя в различных пропорциях их смеси, получим непрерывную кривую C , соединяющую эти точки и проходящую по свойству выпуклости выше и правее отрезка соединяющей их прямой. Так как все точки полученной кривой принадлежат множеству Γ , то и все порождаемые ими прямоугольники, как показано выше, также принадлежат множеству Γ . Отсюда следует, что все точки отрезка, соединяющего исходные точки, также принадлежат множеству Γ . Заметим, что если Q_1 и Q_2 лежат в первом и третьем квадрантах по отношению друг к другу, то этот результат тривиален, тогда отрезок, соединяющий их, лежит внутри прямоугольника, порожденного одной из этих точек.

Изложенного достаточно для того, чтобы из него следовали утверждения, сформулированные в начале этого раздела, а именно, что множество Γ выпукло, совпадает с G_0 и если взять наибольшее достижимое R_{12} и затем для этого R_{12} — наибольшее достижимое R_{21} , то все точки прямоугольника, порожденного точкой (R_{12}, R_{21}) , являются достижимыми. Аналогичный факт справедлив и для наибольшего достижимого R_{21} .

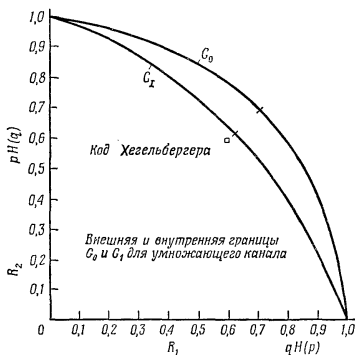
Напомним здесь, что множество точек, достижимых при независимых распределениях $P_1\{x_1, x_2\} = P\{x_1\}P\{x_2\}$, не обязательно является выпуклым множеством. Это показывает пример табл. 1.

Из материала этого раздела следует также, что концевые точки кривых, задающих внутреннюю и внешнюю границы (т. е. точки, в которых они пересекают координатные оси), совпадают. Это следует из того, что, как было показано, наибольшее значение R_{12} может быть получено с использованием только одного фиксированного значения x_2 с вероятностью 1, а тогда $P\{x_1, x_2\}$ сводится к произведению независимых вероятностей.

13. Пример, в котором внутренняя и внешняя границы различны

Внутренняя и внешняя границы для области пропускной способности, которые были описаны выше, не всегда совпадают. Это показал Давид Блекуэлл на примере двоичного умножающего канала,

определяемого соотношением $y_1 = y_2 = x_1 \cdot x_2$. Внутренняя и внешняя границы для этого канала могут быть явно вычислены и их



Р и с. 13.

график представлен на рис. 13. Видно, что они существенно различны, особенно в средней части. Вычисление внутренней границы в этом случае сводится к нахождению огибающей точек

$$\begin{aligned} R_{12} &= -p_2 [p_1 \log p_1 + (1-p_1) \log (1-p_1)], \\ R_{21} &= -p_1 [p_2 \log p_2 + (1-p_2) \log (1-p_2)]. \end{aligned} \quad (45)$$

Это есть скорости при независимых распределениях вероятностей на обоих концах: p_1 — вероятность использования 1 на конце 1, p_2 — вероятность использования 1 на конце 2. Вычисляя эти скорости при различных p_1 и p_2 , можно получить огибающую, изображенную на рисунке.

При вычислении внешней границы необходимо использовать огибающую скоростей для зависимых распределений вероятностей. Однако легко видеть, что любое распределение, у которого вероятность $P\{0,0\}$ положительна, можно улучшить, приписывая эту вероятность одной из других возможных пар. Поэтому снова имеется двухпараметрическое семейство точек (так как сумма остальных

трех вероятностей должна равняться 1). Если эти вероятности обозначить

$$p_1 = P\{1,0\}, \quad p_2 = P\{0,1\}, \quad 1 - p_1 - p_2 = P\{1,1\},$$

то получим скорости

$$R_{12} = -(1-p_1) \left[\frac{p_2}{1-p_1} \log \frac{p_2}{1-p_1} + \left(1 - \frac{p_2}{1-p_1} \right) \log \left(1 - \frac{p_2}{1-p_1} \right) \right], \quad (46)$$

$$R_{21} = -(1-p_2) \left[\frac{p_1}{1-p_2} \log \frac{p_1}{1-p_2} + \left(1 - \frac{p_1}{1-p_2} \right) \log \left(1 - \frac{p_1}{1-p_2} \right) \right].$$

Здесь снова прямое вычисление при различных значениях p_1 и p_2 приводит к огибающей, изображенной на рисунке.

В связи с этим каналом Хегельбергер (D. W. Hagelbarger) предложил интересный и простой код (хотя и не являющийся блоковым), который свободен от ошибок и имеет средние скорости передачи $R_{12} = R_{21} = 0,571$, находящиеся на незначительном расстоянии от нашей нижней границы. Его метод кодирования заключается в следующем: 0 и 1 передаются с каждого конца с независимыми вероятностями $1/2, 1/2$. Если на выходе получается 0, тогда в начале следующего двоичного сигнала посылается сигнал, противоположный только что посланному перед ним сигналу. Этот прием производится на обоих концах. Если же получается 1, тогда на обоих концах переходят к передаче следующего двоичного знака сообщения. Можно заметить, что в среднем $3/4$ времени занимает посылка противоположных сигналов и одна четверть времени употребляется для передачи нового сообщения. Таким образом, среднее число использований канала, нужных для передачи одного двоичного сообщения, составляет $3/4 \cdot 2 + 1/4 \cdot 1 = 7/4$. Средняя скорость передачи по обоим направлениям есть $4/7 = 0,571$. Кроме того, нетрудно видеть, что передаваемые двоичные сообщения могут быть декодированы без ошибок для каждого направления передачи.

Используя источники сообщений на каждом конце канала с несколько смещенными вероятностями, можно несколько улучшить схему Хегельбергера. Так, если 1 встречается как сигнал двоичного сообщения с вероятностью 0,63, а 0 — с вероятностью 0,37, то по обоим направлениям получаются скорости

$$R_{12} = R_{21} = \frac{-0,63 \log 0,63 - 0,37 \log 0,37}{1 - (0,63)^2} = 0,593. \quad (47)$$

В одном из последующих разделов будет получен результат, который в принципе позволяет находить для канала точную область

пропускной способности. Однако при этом используется предельный процесс над словами возрастающей длины и, следовательно, в большинстве случаев вычисления очень трудны. Наоборот, для вычисления верхней и нижней границы требуется только максимизировать некоторые выражения для единственной передаваемой буквы по каждому направлению. Хотя и это иногда требует значительных вычислений, все же для не очень сложных каналов вычисления могут быть доведены до конца.

14. Достижимость внешней границы при зависимых источниках

В связи с рассмотрением внешней границы отметим интересную интерпретацию, относящуюся к несколько более общим системам связи. Предположим, что источники сообщений на обоих концах нашего канала являются не независимыми, а статистически зависимыми. Такой может быть информация о погоде, подлежащая передаче из Бостона в Нью-Йорк и из Нью-Йорка в Бостон. Погода в этих городах, конечно, статистически зависима. Если бы имела место зависимость, приспособленная к каналу или если бы было можно преобразовать сообщение, с тем, чтобы было выполнено вышесказанное, то при передаче можно достигнуть скорости, лежащей на внешней границе. Пусть, например, для только что рассмотренного умножающего канала сообщения на обоих концах состоят из потока двоичных знаков, встречающихся с зависимыми вероятностями, приведенными в табл. III. Последовательные пары x_1, x_2 предполагаются независимыми. Тогда при простой передаче этих потоков

Таблица III

	x_2	
	0	1
x_1	0	0,275
	1	0,45

по каналу (без всякой обработки) кривая внешней границы достигается в ее средней точке.

Неизвестно, возможно ли это в общем случае. Всегда ли существует последовательность пар зависимых источников, которые могут быть закодированы так, чтобы получить скорости R_1, R_2 , отличающиеся меньше чем на ϵ от любой точки внешней границы? Во всяком случае, это утверждение часто бывает справедливым в случае канала без памяти и без шума, т. е. когда y_1 и y_2 являются

точными функциями от x_1 и x_2 . Пара источников, определяемых распределением $P\{x_1, x_2\}$, при котором получается интересующая нас точка, часто оказывается так же, как в приведенном выше примере, пригодной и без дополнительного кодирования.

Внутренняя граница также имеет интересную интерпретацию. Если искусственно ограничить класс кодов только такими, при которых передаваемые последовательности на каждом конце канала зависят только от передаваемого сообщения и не зависят от полученной на этом конце последовательности, то внутренняя граница действительно будет областью пропускной способности. Это вытекает из того, что в таком случае на каждом этапе передачи (т. е. при заданном индексе передаваемой буквы) между двумя следующими передаваемыми буквами нет зависимости. Отсюда следует, что общее векторное изменение неопределенности ограничено суммой n векторов, каждый из которых соответствует независимому распределению вероятностей. Детали этого доказательства предоставляются читателю. Требование независимости сохраняется также в том случае, если передающие и принимающие пункты на каждом конце канала находятся в разных местах и прямой связи не имеют.

15. Общий метод нахождения области пропускной способности двустороннего канала

Для данного двустороннего канала K без памяти определим ряд производных каналов $K_1, K_2, \dots$. Они также будут каналами без памяти и область пропускной способности канала K будет вычисляться как предел внутренних границ для ряда K_n .

Канал K_1 идентичен каналу K . Производный канал K_2 является таким каналом, у которого входные буквы в действительности есть стратегии при работе канала K для блоков из двух входных букв. Таким образом, входные буквы на конце 1 для K_2 состоят из пар $[x_1^1, f(x_1^1, y_1^1)]$. Здесь x_1^1 есть первая передаваемая буква пары, принимающая a значений возможных входных букв канала K . Функция $f(x_1^1, y_1^1)$ представляет собой любую функцию от первой входной буквы x_1^1 и выходной буквы y_1^1 со значением, равным второй входной букве x_2^1 . Таким образом, эта функция может рассматриваться как некоторое правило для выбора второй входной буквы на конце 1 в зависимости от первой входной буквы и полученной первой выходной буквы. Если предположить, что x_1^1 принимает a значений и y_1^1 — b значений, тогда пара (x_1^1, y_1^1) принимает ab значений и, так как функция f может принимать a значений, то всего имеется a^{ab} возможных функций. Таким образом, всего имеется $a \cdot a^{ab}$ возможных пар $[x_1^1, f(x_1^1, y_1^1)]$ или возможных входных букв на конце 1 канала K_2 .

Аналогичным образом на конце 2 рассмотрим пары $[x_2^1, g(x_2^1, y_2^1)]$. Здесь g есть всевозможные функции от первых полученных и переданных букв на конце 2, принимающие значения из алфавита x_2 . Таким образом, этих пар имеется $c \cdot c^d$, где c и d являются объемами входного и выходного алфавитов на конце 2.

Пары $[x_1^1, f(x_1^1, y_1^1)]$ и $[x_2^1, g(x_2^1, y_2^1)]$ могут рассматриваться как стратегии при использовании канала K для передачи последовательностей из двух букв, если вторая буква может зависеть от первой передаваемой и первой полученной буквы. Методика здесь очень похожа на методику, используемую в теории игр. Последовательность ходов игрока (для которого имеющаяся у него информация для выбора хода возрастает по мере увеличения номера хода) заменяется одним ходом, которым он выбирает всю стратегию. Стратегия описывает, что должен делать игрок на каждом этапе во всех возможных обстоятельствах. Таким образом, многоходовая игра сводится к одноходовой игре с ходом, выбираемым из большой совокупности.

Выходными буквами для K_2 на конце 1 являются пары (y_1^1, y_2^1) и на конце 2 пары (y_1^2, y_2^2) , т. е. пары полученных букв на обоих концах. Переходные вероятности для K_2 есть вероятности того, что появятся данные выходные пары при условии, что при вводе в канал K фиксированной пары букв были использованы данные стратегии. Таким образом,

$$P_{K_2}\{(y_1^1, y_2^1), (y_2^1, y_2^2) | [x_1^1, f(x_1^1, y_1^1)], [x_2^1, g(x_2^1, y_2^1)]\} = \\ = P_K\{y_1^1, y_2^1 | x_1^1, x_2^1\} P_K\{y_2^2, y_2^2 | f(x_1^1, y_1^1), g(x_2^1, y_2^1)\}. \quad (48)$$

Аналогичным путем определяются каналы $K_3, K_4, \dots$. Таким образом, канал K_n может рассматриваться как канал, соответствующий n -кратному использованию канала K ; последовательности входных букв на каком-либо его конце являются функциями от предыдущих входных и выходных букв на этом конце. Поэтому входные буквы на конце 1 являются n -мерными векторами

$$[x_1^1, f_1(x_1^1, y_1^1), \dots, f_{n-1}(x_1^1, x_2^1, \dots, x_1^{n-1}, y_1^1, y_2^1, \dots, y_1^{n-1})] \quad (49)$$

из алфавита c

$$aa^{ab} a^{(ab)^2} \dots a^{(ab)^{n-1}} = a^{[(ab)^n - 1]/ab - 1} \quad (50)$$

буквами. Выходными буквами на конце 1 являются n -мерные векторы

$$(y_1^1, y_2^1, \dots, y_n^1), \quad (51)$$

принимающие поэтому значения из алфавита с b^n обобщенными буквами. Переходные вероятности для канала K_n выражаются через

переходные вероятности для канала K с помощью соотношения, обобщающего соотношение (39)

$$P_{K_n}\{y_1^1, y_1^2, \dots, y_1^n | (x_1^1, f_1, f_2, \dots, f_{n-1}), (x_2^1, g_1, g_2, \dots, g_{n-1})\} = (52) \\ = \prod_{i=1}^n P_K\{y_1^i | f_{i-1}, g_{i-1}\}.$$

Канал K_n может рассматриваться, следовательно, как канал без памяти, свойства которого идентичны свойствам канала K при передаче по нему целых блоков длины n , причем допускается влияние передаваемых и получаемых букв внутри блока на последующий выбор. Для каждого канала K_n в принципе можно вычислить нижнюю границу области пропускной способности. Нижнюю границу для канала K_n надо при сравнении с границей для канала K умножить на множитель $1/n$, так как K_n соответствует n -кратному использованию канала K .

Теорема 5. Пусть B_n есть нижняя граница для области пропускной способности производного канала K_n , растянутая в $1/n$ раз. Тогда при $n \rightarrow \infty$ области B_n стремятся к предельной области B , которая содержит каждую из них и является областью пропускной способности для канала K .

Доказательство. Сначала докажем прямое утверждение теоремы: если (R_{12}, R_{21}) есть любая точка из некоторого B_n и ε — любое положительное число, то можно построить блочный код со скоростями передачи, не меньшими $R_{12} - \varepsilon$ и $R_{21} - \varepsilon$ по обоим направлениям соответственно и при вероятности ошибки $P_e < \varepsilon$. Это немедленно следует из предыдущих результатов, если производный канал K_n и его связь с внутренней границей B_n определены надлежащим образом. K_n есть некоторый канал без памяти, и по теореме 3 можно найти коды со скоростями передачи, как угодно близкими к скоростям R_{12} , R_{21} , из B_n при произвольно малой вероятности ошибки. Эти коды являются последовательностями букв из алфавита K_n . Они соответствуют, таким образом, последовательностям стратегий для исходного канала K при передаче блоков длины n .

Таким образом, эти коды могут быть непосредственно преобразованы в коды для канала K с длиной в n раз большей и с сохранением всех статистических свойств в том числе и вероятности ошибки. Эти коды тогда могут быть интерпретированы как коды для канала K со скоростями передачи, в n раз большими, при той же вероятности ошибки. Действительно, из теоремы 3 следует, что для любой пары скоростей, находящихся строго внутри B_n , можно найти коды с вероятностями ошибок, убывающими по крайней мере экспоненциально с ростом длины кода.

Докажем теперь, что при росте n области B_n стремятся к предельной области B , которая содержит все частные B_n . Под пре-

дельной областью понимается множество B таких точек, что для любой точки P из B и произвольного $\varepsilon > 0$ существует такое n_0 , что для всех $n > n_0$ существуют точки из B_n , отстоящие от точки P менее чем на ε ; в то же время для любой точки P , не принадлежащей B , существуют такие ε и n_0 , что для любого $n > n_0$ в B_n не содержится точек, отстоящих от P менее чем на ε . Во-первых, B_n содержится в B_{kn} для любого k . Это вытекает из того, что стратегии для B_{kn} включают в себя как частный случай стратегии, зависящие только от подблоков длины n . Поэтому все достижимые точки для канала K при независимых распределениях являются достижимыми также и для K_{kn} и выпуклая оболочка последнего множества должна содержать в себе выпуклую оболочку первого множества.

Отсюда следует, что множества B_{kn} стремятся к пределу B , являющемуся объединением всех B_{kn} и их предельных точек. Множество B содержит также и B_{n_1} при любом n_1 . Ведь если, например, n и n_1 — общее кратное, например, nn_1 , то B содержит B_{nn_1} , в то время как B_{nn_1} содержит B_{n_1} .

Кроме того, любая достижимая точка для канала K_{kn} может быть достигнута и в канале $K_{kn+\alpha}$ для $0 \leq \alpha \leq n$ умножением обеих координат на коэффициент, не больший чем $k/(k+1)$. Это следует из того, что можно использовать стратегии для K_{kn} с приписываемыми вслед за ними последовательностями по α первых букв из алфавитов x_1 и x_2 . (То есть распространение распределения до длины $kn+\alpha$ происходит за счет приписывания «пустых» букв.) Единственное различие состоит в нормирующем множителе, равном $1/(\text{длина блока})$. При достаточно больших k отличие этого множителя от 1, равное $1/k+1$, может быть сделано как угодно малым. Таким образом, для любого $\varepsilon > 0$ и любой точки P из B для всех достаточно больших n_1 существует точка из B_{n_1} , отстоящая от P менее чем на ε .

При рассмотрении обратного утверждения теоремы предположим, что имеется блочный код длины n со скоростями передачи (R_1, R_2) , которые соответствуют точке, внешней для множества B , кратчайшее расстояние от которой до B равно ε . Так как B содержит B_n , то кратчайшее расстояние от этой точки до B_n не менее ε . Можно рассматривать этот код как блочный код длины 1 для канала K_n , а тогда сообщения m_1 и m_2 отображаются непосредственно во «входные буквы» канала K_n без функциональной зависимости от полученных букв. Тогда так как m_1 и m_2 независимы, то распределения вероятностей для наших входных букв являются независимыми, а этого достаточно, чтобы внешняя и внутренняя границы совпадали. Итак, при рассмотренном коде вероятность ошибки ограничена снизу некоторой величиной, большей нуля и зависящей от ε , но не зависящей от n .

16. Двусторонние каналы с памятью

Общий дискретный двусторонний канал с памятью определяется совокупностью условных вероятностей

$$P\{y_{1n}, y_{2n}/x_{11}, x_{12}, \dots, x_{1n}; x_{21}, x_{22}, \dots, x_{2n}; \\ y_{11}, y_{12}, \dots, y_{1n-1}; y_{21}, y_{22}, \dots, y_{2n-1}\}. \quad (53)$$

Это есть вероятность n -й выходной пары y_{1n}, y_{2n} при условии что известно «прошлое» с момента $t=0$, т. е. входные и выходные последовательности с начала работы канала. В таком общем случае эти вероятности могут изменяться совершенно произвольным образом при росте n . Если не налагать никаких дальнейших ограничений, то этот случай является слишком общим, чтобы быть полезным или интересным. Поэтому необходимо дополнительно наложить некоторые ограничения разумной общности, которые, однако, должны обеспечивать некоторую устойчивость поведения канала и поэтому выполнение важных теорем кодирования. Например, можно требовать ограничения влияния прошлого, так что вероятности букв будут тогда зависеть только от ограниченного отрезка прошлого. (Если известны последние α входных и выходных букв, то более ранние входные и выходные буквы не оказывают влияния на условные вероятности.) Будем, однако, использовать некоторое условие, которое, вообще говоря, является более общим, а также более правдоподобным в реальных приложениях.

Будем говорить, что двусторонний канал обладает *свойством возвратности состояния*, если выполняется следующее условие. Существует число d , такое, что для любых входных и выходных последовательностей длины n , $X_{1n}, X_{2n}, Y_{1n}, Y_{2n}$ существуют две функции $f(X_{1n}, Y_{1n}), g(X_{2n}, Y_{2n})$, значениями которых являются последовательности входных букв той же самой длины, меньшей d , и такие, что если эти последовательности f и g передать по каналу, то он вернется в первоначальное состояние¹⁾. Таким образом, условные вероятности после этого становятся такими же, как если бы канал начинал свою работу с нулевого момента времени.

Свойство возвратности состояния встречается в реальных физических системах связи, где часто имеется «нулевой» вход, такой, что при применении его в течение достаточно большого времени можно исключить влияние прошлого. Заметим также, что свойство возвратности состояния может встречаться даже в каналах с бесконечным множеством внутренних состояний, если только можно возвратиться в «основное» состояние через ограниченное число шагов.

¹⁾ Математически строгое введение понятия состояния канала можно найти в работе: Блекуэлл, Брейман, Томасян, сб. *Математика*, 4:5 (1960), ИЛ, 123. — *Прим. ред.*

Смысл условия возвратности состояния состоит в том, что, имея блочный код для такого канала, можно применить к входным буквам этого кода функции f и g на обоих концах и затем повторно использовать этот код. Таким образом, если такой код имеет длину n и при однократном его использовании скорости передачи равны R_1 и R_2 , а вероятности ошибок P_{e1} и P_{e2} , можно непрерывно вести передачу со скоростями $R'_1 \geq nR_1/(n+d)$ и $R'_2 \geq nR_2/(n+d)$ при вероятностях ошибок $P'_{e1} \leq P_{e1}$ и $P'_{e2} \leq P_{e2}$.

Для канала с возвратным состоянием можно рассмотреть стратегии для первых n букв точно так же, как это было в случае канала без памяти, и найти соответствующую внутреннюю границу B_n для области пропускной способности (в масштабе $1/n$).

Определим область B , которую можно назвать верхним пределом областей B_n . А именно B состоит из всевозможных точек, которые принадлежат бесконечному числу B_n , и из предельных точек этого множества.

Теорема 6. Пусть (R_1, R_2) есть любая точка из области B . Пусть n_0 — любое число, а ε_1 и ε_2 — любые положительные числа. Тогда существует блочный код длины $n > n_0$ со скоростями передач R'_1, R'_2 , такими, что $|R_1 - R'_1| < \varepsilon_1, |R_2 - R'_2| < \varepsilon_1$ и вероятности ошибок $P_{e1} < \varepsilon_2, P_{e2} < \varepsilon_2$. Наоборот, если (R_1, R_2) не принадлежит B , то существуют такие n_0 и $\delta > 0$, что для любого блочного кода длиной, большей n_0 , либо $P_{e1} > \delta$, либо $P_{e2} > \delta$ (либо выполняются оба неравенства одновременно).

Чтобы доказать первую часть теоремы, выберем некоторое достаточно большое $n_1 > n_0$, так, чтобы обе величины $dR_1/(d+n)$ и $dR_2/(d+n)$ были меньше $\varepsilon_1/2$. Так как точка (R_1, R_2) содержится в бесконечной последовательности B_n , то это возможно сделать. Построим теперь блочный код на основании n_1 -кратного использования канала при передаче отдельных «букв» со скоростями, отличающимися от (R_1, R_2) менее чем на $\varepsilon_1/2$ при вероятностях ошибок, меньших ε_2 . К каждой из «букв» этого кода применим функции, которые возвращают канал в первоначальное состояние. Таким образом, получим коды при произвольно малых вероятностях ошибок, меньших ε_2 , аппроксимирующие скорости R_1, R_2 и имеющие сколь угодно большую длину блоков.

Чтобы доказать обратное утверждение, предположим, что (R_1, R_2) не принадлежит B . Тогда для некоторого n_0 любое B_n с номером $n > n_0$ лежит вне круга некоторого радиуса, скажем ε_2 , с центром в точке (R_1, R_2) , так как иначе (R_1, R_2) было бы предельной точкой множества B_n . Предположим, что имеется код длины $n_1 > n_0$. Тогда вероятности ошибок при его использовании ограничены снизу величиной, большей нуля, так как снова имеется случай, где получается независимость между «буквами».

Область B может быть названа областью пропускной способности для такого канала с возвратным состоянием. Легко показать, что B обладает теми же свойствами выпуклости, как и область пропускной способности G для канала без памяти. Конечно, практическое вычисление B для конкретных каналов является даже менее реальной задачей, чем в случае каналов без памяти.

17. Обобщение на каналы с T концами

Многие методы и приемы, развитые выше, могут быть обобщены на каналы с тремя и более концами. Однако в этих более сложных случаях появляются некоторые определенно новые явления. В другой работе будет рассмотрен случай канала с двумя или более концами, имеющими только входы, и с одним концом, имеющим только выход; в этом случае может быть найдено легкое и простое решение задачи об области пропускной способности.

Приложение

Границы для вероятностей ошибок в терминах производящих функций для моментов

Предположим, что вероятности $P\{x_1\}$ приписаны входным буквам на конце 1 и $P\{x_2\}$ — входным буквам на конце 2. (Заметим, что здесь речь идет о буквах, а не о словах, как в теореме 2.) Можно тогда вычислить логарифм от производящей функции для моментов взаимной информации между входными буквами на конце 1 и входными и выходными парами букв на конце 2. (Это есть логарифм производящей функции для моментов распределения q_{12} , когда $n=1$.) Выражения для этой величины и ей подобной для второго направления таковы:

$$\begin{aligned} \mu_1(s) &= \log \sum_{x_1, x_2, y_2} P\{x_1, x_2, y_2\} \exp\left(s \log \frac{P\{x_1, x_2, y_2\}}{P\{x_1\}P\{x_2, y_2\}}\right) = \\ &= \lg \sum_{x_1, x_2, y_2} \frac{P\{x_1, x_2, y_2\}^{s+1}}{P\{x_1\}^s P\{x_2, y_2\}^s}, \end{aligned} \quad (54)$$

$$\mu_2(s) = \log \sum_{x_1, x_2, y_1} \frac{P\{x_1, x_2, y_1\}^{s+1}}{P\{x_2\}^s P\{x_1, y_1\}^s}. \quad (55)$$

Эти функции μ_1 и μ_2 могут быть использованы при ограничении «хвостов» распределений q_{12} и q_{21} , полученных при сложении вместе n одинаково распределенных выборок. Действительно, Чернов¹⁾ показал, что хвост слева от среднего может быть ограничен

¹⁾ Chernoff H., A measure of asymptotic efficiency for tests of a hypothesis on the sum of observations, *Ann. Math. Statist.*, 23 (1952), 493.

следующим образом:

$$\begin{aligned} Q_{12} [n\mu'_1(s_1)] &\leq \exp \{n [\mu_1(s_1) - s_1\mu'_1(s_1)]\}, \quad s_1 \leq 0, \\ Q_{12} [n\mu'_2(s_2)] &\leq \exp \{n [\mu_2(s_2) - s_2\mu'_2(s_2)]\}, \quad s_2 \leq 0. \end{aligned} \quad (56)$$

Таким образом, выбирая произвольное отрицательное s_1 , по этим формулам получим границу для значения функции распределения в точке $n\mu'_1(s_1)$. Можно показать, что $\mu'(s)$ является монотонно возрастающей функцией и что $\mu'(0)$ есть среднее ее распределения. Минимум $\mu'(s)$ соответствует минимально возможному значению рассматриваемой случайной величины, в данном случае минимуму $I(x_1; x_2, y_2)$. Таким образом, можно найти такое значение s_1 , чтобы $\mu_1(s_1)$ находилось в любом месте между $I_{\min}(x_1; x_2, y_2)$ и $E(I)$. Конечно, левее $I_{\min}$ функция распределения тождественно равна 0, а правее $E(I)$ функция распределения стремится к 1 при возрастании n .

Желательно воспользоваться этими результатами, чтобы получить более точные границы для P_{e1} и P_{e2} , фигурирующих в теореме 2. Вспоминая, что в этой теореме θ_1 и θ_2 являются произвольными, постараемся выбрать их так, чтобы экспоненциальные границы для обоих членов совпадали. Это хороший способ выбора θ_1 и θ_2 , если требуется поддерживать общую границу как можно меньшей. Первый член ограничен выражением $\exp \{n [\mu_1(s_1) - s_1\mu'_1(s_1)]\}$, где s_1 таково, что $\mu'_1(s_1) = R_1 + \theta_1$, второй член равен $\exp(-n\theta_1)$. Объединяя эти равенства, получим

$$\mu_1(s_1) - s_1\mu'_1(s_1) = -\theta_1, \quad R_1 + \theta_1 = \mu'_1(s_1). \quad (57)$$

Исключая θ_1 , будем иметь

$$R_1 = \mu_1(s_1) - (s_1 - 1)\mu'_1(s_1) \quad (58)$$

и

$$E(P_{e1}) \leq 2 \exp \{n [\mu_1(s_1) - s_1\mu'_1(s_1)]\}. \quad (59)$$

Это следует из того, что оба члена теперь равны, и каждый из них мажорируется величиной $\exp \{n [\mu_1(s_1) - s_1\mu'_1(s_1)]\}$. Аналогично для

$$R_2 = \mu_2(s_2) - (s_2 - 1)\mu'_2(s_2) \quad (60)$$

имеем

$$E(P_{e2}) \leq 2 \exp \{n [\mu_2(s_2) - s_2\mu'_2(s_2)]\}. \quad (61)$$

Эти выражения могут быть названы параметрическими границами, выраженными через параметры s_1 и s_2 . Надо выбрать s_1 и s_2 так, чтобы скорости R_1 и R_2 приняли желаемые значения. Если под-

ставить эти значения s_1 и s_2 в остальные формулы, то получим границы для вероятностей ошибок.

Производная от R_1 по s_1 равна $-(s_1 - 1)\mu_1''(s_1)$, т. е. величине, являющейся всегда положительной при отрицательном s_1 , за исключением особого случая, когда $\mu''(0) = 0$. Таким образом, R_1 является монотонно возрастающей функцией от s_1 , принимающей значения от $-I_{\min} - \log P\{I_{\min}\}$ до $E(I)$ при изменении s_1 от $-\infty$ до 0. Тем временем величина, стоящая в квадратных скобках в выражении $E(P_{e1})$, т. е. $\mu_1(s_1) - s_1\mu_1'(s_1)$, изменяется от $\log P\{I_{\min}\}$ до нуля. Скорость, соответствующая значению $s_1 = -\infty$, т. е. величине $-I_{\min} - \log P\{I_{\min}\}$, может быть как положительной, так и отрицательной. Если она отрицательна (или равна 0), то покрывается целый промежуток скоростей от 0 до $E(I)$. Однако если она положительна, то имеется пробел для скоростей от $R_1 = 0$ до этой конечной точки. Это означает, что уравнение, приводящее к двум равным экспоненциальным членам, не имеет решения при скоростях в этом интервале. В таком случае, чтобы получить хорошую границу, лучше всего выбрать θ_1 так, чтобы $n(R_1 + \theta_1)$ было меньше $I_{\min}$, скажем, равным $I_{\min} - \varepsilon$. Тогда $q_{12}[n(R_1 + \theta_1)] = 0$ и остается только второй член $\exp(\theta_1 n)$. Таким образом, $\exp[-n(I_{\min} - R_1 - \varepsilon)]$ является границей для P_e . Это справедливо для любого $\varepsilon > 0$. Так как можно построить такие коды для любого положительного ε и так как имеется лишь конечное число кодов, то отсюда вытекает, что может быть построен код, удовлетворяющий этому неравенству при $\varepsilon = 0$. Таким образом, можно сказать, что

$$E(P_{e1}) \leq \exp[-n(I_{\min} - R_1)], \quad R_1 \leq I_{\min}. \quad (62)$$

Конечно, в точности аналогичное утверждение справедливо и для кодирования по обратному направлению. Комбинируя и суммируя изложенное, получим следующую теорему.

Теорема 7. Пусть для двустороннего канала K без памяти множества $P\{x_1\}$ и $P\{x_2\}$ являются распределением вероятностей на входных алфавитах и предположим, что при этом логарифмы от производящей функции для моментов взаимной информации $\mu_1(s_1)$ и $\mu_2(s_2)$ равны

$$\begin{aligned} \mu_1(s_1) &= \log \sum_{x_1, x_2, y_2} \frac{P\{x_1, x_2, y_2\}^{s_1+1}}{P\{x_1\}^{s_1} P\{x_2, y_2\}^{s_1}}, \\ \mu_2(s_2) &= \log \sum_{x_1, x_2, y_2} \frac{P\{x_1, x_2, y_1\}^{s_2+1}}{P\{x_2\}^{s_2} P\{x_1, y_1\}^{s_2}}. \end{aligned} \quad (63)$$

Пусть $M_1 = \exp(R_1 n)$, $M_2 = \exp(R_2 n)$ и пусть s_1 и s_2 являются решениями (когда они существуют) уравнений

$$\begin{aligned} R_1 &= \mu_1(s_1) - (s_1 + 1) \mu'_1(s_1), \\ R_2 &= \mu_2(s_2) - (s_2 + 1) \mu'_2(s_2). \end{aligned} \quad (64)$$

Решение s_1 существует, если

$$-I_{\min}(x_1; x_2, y_2) - \log P\{I_{\min}(x_1; x_2, y_2)\} \leq R_1 \leq E[I(x_1; x_2, y_2)], \quad (65)$$

аналогично и для s_2 . Если оба решения s_1 и s_2 существуют, то существует некоторая пара кодов длины n для канала K и передачи M_1 и M_2 сообщений по обоим направлениям соответственно с вероятностями ошибок, удовлетворяющим условиям

$$\begin{aligned} P_{e1} &\leq 4 \exp\{n[\mu_1(s_1) - s_1 \mu'_1(s_1)]\}, \\ P_{e2} &\leq 4 \exp\{n[\mu_2(s_2) - s_2 \mu'_2(s_2)]\}. \end{aligned} \quad (66)$$

Если одно из R (или оба) так малы, что соответствующие s не существуют, то существует пара кодов с вероятностями ошибок, ограниченных соответственно выражениями

$$P_{e1} \leq 2 \exp\{-n[I(x_1; x_2, y_2) - R_1]\} \quad (67)$$

и

$$P_{e2} \leq 2 \exp\{-n[I(x_2; x_1, y_1) - R_2]\}. \quad (68)$$

Таким образом, если существует s_1 и не существует s_2 , то используются неравенства (66). Если ни одно из них не существует, то справедливы неравенства (67) и (68).

Разное

За последние несколько лет теория информации превратилась в своего рода бандвагон от науки. Появившись на свет в качестве специального метода в теории связи, она заняла выдающееся место как в популярной, так и в научной литературе. Это можно объяснить отчасти ее связью с такими модными областями науки и техники, как кибернетика, теория автоматов, теория вычислительных машин, а отчасти новизной ее тематики. В результате всего этого значение теории информации было, возможно, преувеличено и раздуто до пределов, превышающих ее реальные достижения. Ученые различных специальностей, привлеченные поднятым шумом и перспективами новых направлений исследования, используют идеи теории информации при решении своих частных задач. Так, теория информации нашла применение в биологии, психологии, лингвистике, теоретической физике, экономике, теории организации производства и во многих других областях науки и техники. Короче говоря, сейчас теория информации, как модный опьяняющий напиток, кружит голову всем вокруг.

Для всех, кто работает в области теории информации, такая широкая популярность несомненно приятна и стимулирует их работу, но такая популярность в то же время и настораживает. Сознывая, что теория информации является сильным средством решения проблем теории связи (и в этом отношении ее значение будет возрастать), нельзя забывать, что она не является панацеей для инженера-связиста и тем более для представителей всех других специальностей. Очень редко удается открыть одновременно

¹⁾ Shannon C., The Bandwagon, *Trans. IRE*, IT-2, № 1 (1956), 3. Слово бандвагон (bandwagon) в Америке означает политическую партию, добившуюся популярности и победившую на выборах, или просто группу лиц, программа которых находит широкую поддержку населения. Слово состоит из двух частей: «band» (оркестр, джаз) и «wagon» (повозка, карета) и, возможно, связано с существовавшим обычаем, по которому победивший на выборах кандидат проезжал по городу в открытой машине с джазом. Но, кроме того, слово «band» применяется в теории связи, где оно означает полосу пропускания частот, так что в этом заглавии есть некоторая игра слов.— *Прим. перев.*

несколько тайн природы одним и тем же ключом. Знание нашего несколько искусственно созданного благополучия слишком легко может рухнуть, как только в один прекрасный день окажется, что при помощи нескольких магических слов, таких, как *информация, энтропия, избыточность...*, нельзя решить всех нерешенных проблем.

Что можно сделать, чтобы внести в сложившуюся ситуацию ноту умеренности? Во-первых, представителям различных наук следует ясно понимать, что основные положения теории информации касаются очень специфического направления исследования, направления, которое совершенно не обязательно должно оказаться плодотворным в психологии, экономике и в других социальных науках. В самом деле, основу теории информации составляет одна из ветвей математики, т. е. строго дедуктивная система. Поэтому глубокое понимание математической стороны теории информации и ее практических приложений к вопросам общей теории связи является обязательным условием использования теории информации в других областях науки. Я лично полагаю, что многие положения теории информации могут оказаться очень полезными в этих науках; действительно, в ней уже достигнуты некоторые весьма значительные результаты. Однако поиск путей применения теории информации в других областях не сводится к тривиальному переносу терминов из одной области науки в другую. Этот поиск осуществляется в длительном процессе выдвижения новых гипотез и их экспериментальной проверки. Если, например, человек в определенной ситуации ведет себя подобно идеальному декодирующему устройству, то это является экспериментальным фактом, а не математическим выводом и, следовательно, требуется экспериментальная проверка такого поведения на широком фоне различных ситуаций.

Во-вторых, мы должны поддерживать образцовый порядок в своем собственном доме. На понятия теории информации очень большой, даже, может быть, слишком большой спрос. Поэтому мы сейчас должны обратить особое внимание на то, чтобы исследовательская работа в нашей области велась на самом высоком научном уровне, который только возможно обеспечить. Больше исследовать и меньше демонстрировать свои достижения, повысить требовательность к себе — вот что должно быть сейчас нашим лейтмотивом. Исследователям следует публиковать результаты только своих наиболее ценных работ и то лишь после серьезной критики как со своей стороны, так и со стороны своих коллег. Лучше иметь небольшое количество первоклассных статей, чем много слабо продуманных или недоработанных публикаций, которые не принесут чести их авторам и только отнимут время у читателей. Только последовательно придерживаясь строго научной линии, мы сможем достичь реальных успехов в теории связи и укрепить свои позиции.

1. Введение

В опубликованной ранее работе²⁾ были введены понятия энтропии и избыточности языка. Энтропия есть статистический параметр, который измеряет в известном смысле среднее количество информации, приходящейся на одну букву языкового текста. Если данный язык перевести на язык двоичных знаков (0 или 1) наиболее эффективным образом, то энтропия H равна среднему числу двоичных знаков (бит), приходящихся на одну букву исходного языка. Избыточность в свою очередь измеряет количество ограничений в языковом тексте, определяемое его статистической структурой; например, в английском языке наибольшая частота появления буквы E, частое явление буквы H вслед за T и U вслед за Q. По ранее произведенной оценке при учете статистических связей не более чем между восемью последовательными буквами оказалось, что энтропия равна примерно 2,3 бита на одну букву и избыточность составляет около 50%.

С тех пор был найден новый метод для оценки этих количеств, более тонкий и учитывающий длительные статистические связи, влияние отдельных фраз друг на друга и т. д. Этот метод основан на изучении возможности предсказания английского текста: насколько точно может быть предсказана очередная буква, когда известны предыдущие N букв текста. Далее будут приведены результаты некоторых экспериментов по предсказанию и теоретический анализ идеального предсказания. Комбинируя экспериментальные и теоретические результаты, можно дать оценки сверху и снизу для энтропии и избыточности. Из проводимого анализа вытекает, что в нормативном английском литературном тексте длительные статистические связи (до 100 букв) уменьшают энтропию приблизительно на один бит на букву с соответствующей избыточностью в 75%. Избыточность может быть еще выше, если учитывать связи

¹⁾ Shannon C., Prediction and entropy of printed English, *BSTJ*, № 1 (1951), 50.

²⁾ Shannon C., Mathematical theory of communication. (Русский перевод см. стр. 243 данного сборника.— Прим. ред.)

между разными параграфами, главами и т. д. Однако при увеличении длин рассматриваемых текстов рассматриваемые параметры становятся более неустойчивыми и неопределенными и делаются существенно зависимыми от типа изучаемого текста.

2. Вычисление энтропии по статистике английского языка

По одному из методов вычисления энтропии задается ряд последовательных приближений $F_0, F_1, F_2, \dots, F_n$ к H , как к пределу, которые учитывают все большее число и более тонкие статистические закономерности языка. Приближение F_N может быть названо N -граммной энтропией; она измеряет количество информации, или энтропию, с учетом статистических связей не длиннее, чем на N следующих друг за другом букв текста. F_N дается формулой ¹⁾

$$F_N = - \sum_{i,j} p(b_i, j) \log_2 p(b_i, j) = - \sum_{i,j} p(b_i, j) \log_2 p(b_i, j) + \sum p(b_i) \log p(b_i), \quad (1)$$

в которой b_i — блок из $N-1$ буквы $[(N-1)$ -грамма],

j — произвольная буква, следующая за b_i ,

$p(b_i, j)$ — вероятность N -граммы b_i, j ,

$p(b_i, j)$ — условная вероятность буквы j следовать за блоком

b_i , равная $\frac{p(b_i, j)}{p(b_i)}$.

Соотношение (1) можно интерпретировать как формулу для вычисления средней неопределенности (условной энтропии) последующей буквы j , когда известны предыдущие $N-1$ букв. При возрастании N в величине F_N учитываются все более и более далекие статистические связи и энтропия H является предельным значением F_N при $N \rightarrow \infty$,

$$H = \lim_{N \rightarrow \infty} F_N. \quad (2)$$

N -граммная энтропия F_N для малых значений может быть подсчитана из обычных частотных таблиц отдельных букв, двухбуквенных (диграмм) и трехбуквенных сочетаний (триграмм). Если промежутком между буквами и знаками препинания пренебречь, то получим 26-буквенный алфавит и F_0 может быть взята (по определению) равной $\log_2(26)$ или 4,7 бита на букву. F_1 при использовании частоты появления отдельных букв равна

$$F_1 = - \sum_{i=1}^{26} p(i) \log p(i) = 4,14 \text{ битов.} \quad (3)$$

¹⁾ Pratt F., Secret and Urgent, Blue Ribbon Books, 1942.

Диграммное приближение F_2 дает результат

$$F_2 = - \sum_{i,j} p(i, j) \log_2 p_{ij}(j) = - \sum p(i, j) \log_2 p(i, j) + \\ + \sum_i p(i) \log_2 p(i) = 7,70 - 4,14 = 3,56 \text{ битов.} \quad (4)$$

Триграммная энтропия равна

$$F_3 = - \sum p(i, j, k) \log_2 p_{ij}(k) = \\ = - \sum p(i, j, k) \log_2 (i, j, k) + \sum p_{ij} \log_2 p(i, j) = 11,0 - 7,7 = 3,3.$$

Таблицы триграмм, использованные в этих вычислениях, не принимали в расчет трехбуквенных сочетаний, связывающих два слова, к примеру, WOW и OWO в словосочетании TWO WORDS. Для частичной компенсации этого упущения были составлены исправленные таблицы вероятностей трехбуквенных сочетаний $p(i, j, k)$, полученных из вероятностей $p'(i, j, k)$, взятых из таблиц, с помощью следующей грубой формулы:

$$p(i, j, k) = \frac{2,5}{4,5} p'(i, j, k) + \frac{1}{4,5} r(i) p(j, k) + \frac{1}{4,5} p(i, j) s(k),$$

где $r(i)$ есть вероятность того, что буква i находится на последнем месте в слове, а $s(k)$ — вероятность того, что буква k является начальной буквой слова. Таким образом, триграммы внутри слов (в среднем 2,5 на слово) вычислялись в соответствии с таблицей; триграммы, встречающиеся между словами (по одной каждого типа на слово), вычислялись приближенно в предположении независимости последней буквы слова и начальной диграммы следующего слова и, наоборот, последней диграммы слова и начальной буквы следующего слова. В результате этих приближений, а также вследствие того факта, что выборочная ошибка при отождествлении вероятности с выборочной частотой является в этом случае более существенной, полученную величину для F_3 следует считать менее надежной, чем предыдущие.

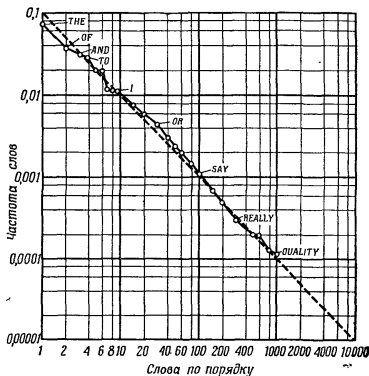
Поскольку таблиц N -граммных частот при $N > 3$ не существует, F_4 , F_5 и т. д. нельзя вычислить тем же путем. Однако были составлены таблицы частоты появления слов¹⁾, и их можно использовать для получения дальнейших приближений. Рис. 1 изображает в логарифмическом масштабе вероятность слов в порядке убывания частоты их появления. Наиболее часто встречающееся английское слово «the» имеет вероятность 0,071 и изображено над 1. Следующее по частоте слово «of» имеет вероятность 0,034 и изображено над 2 и т. д. При использовании логарифмического масштаба как для

¹⁾ Dewey G., Relative frequency of English speech sounds, Harvard University Press, 1923.

вероятности, так и для порядкового номера слова, получающаяся кривая выглядит приблизительно прямой линией с наклоном -1 ; таким образом, если p_n есть вероятность n -го по порядку слова, имеем приблизительно

$$p_n = \frac{0,1}{n}. \quad (6)$$

Ципф¹⁾ указал, что формула вида $p_n = k/n$ дает довольно хорошее приближение вероятностей слов во многих языках. Формула (6),



Р и с. 1.

очевидно, не справедлива при n , стремящемся к ∞ , поскольку общая вероятность $\sum p_n$ должна быть равна единице, в то время как $\sum_{n=1}^{\infty} \frac{0,1}{n}$ равна бесконечности. Если предположить (ввиду отсутствия любой более удовлетворительной оценки), что формула $p_n = \frac{0,1}{n}$ выполняется до тех значений n , пока общая вероятность не станет

¹⁾ Zipf C. V., Human behavior and the principle of least effort, Addison-Wesley Press, 1949.

равной единице, и что $p_n=0$ для больших значений n , то найдем, что наибольшее n равно 8,727. Тогда энтропия равна

$$- \sum_1^{8,727} p_n \log_2 p_n = 11,82 \text{ битов} \quad (7)$$

или $11,82/4,5=2,62$ бита на букву, поскольку средняя длина слова в английском языке равна 4,5 буквы. Можно попытаться отождествить это значение с $F_{4,5}$, но в действительности ордината кривой F_N при $N=4,5$ будет лежать над этим значением. Причина этого заключается в том, что в F_4 и F_5 учтены группы из четырех или пяти букв независимо от подразделения на слова. Слово является связанной группой букв с сильными внутренними статистическими связями и, следовательно, N -граммы, находящиеся внутри слов, несут в себе большие статистические ограничения, чем N -граммы, включающие промежутки между словами. В результате этого полученная нами оценка 2,62 бита на букву более близка, скажем, к F_5 или F_6 .

Аналогичные вычисления были проделаны с учетом пробела между словами как дополнительной буквы, приводящей к алфавиту из 27 букв. Ниже собраны результаты для 26-ти и 27-ми буквенных вычислений

	F_0	F_1	F_2	F_3	$F_{\text{слово}}$
26 букв	4,70	4,14	3,56	3,3	2,62
27 букв	4,76	4,03	3,32	3,1	2,14

Оценка 2,3 для F_8 , упомянутая выше, была найдена несколькими методами, один из которых состоит в экстраполяции до F_8 приведенных результатов для 26-буквенного алфавита. Поскольку пробел между словами является почти полностью избыточным, когда рассматриваются последовательности из многих слов, то значения F_N для 27-буквенного алфавита равны 4,5/5,5 или 0,818 от F_N для 26-буквенного алфавита при больших N .

3. Предсказание английского текста

В новом методе оценки энтропии английского языка используется тот факт, что каждый, говорящий на этом языке, обладает огромными трудно учитываемыми сведениями о статистике языка. Знакомство со словами, идиомами, стандартными оборотами и грамматикой позволяет выправить неправильные или пополнить пропущенные буквы при чтении корректур или дополнить неоконченную фразу в разговоре. Экспериментальная демонстрация степени возможности предсказания английского текста может быть следующей: выберем короткий отрывок текста, не известный отгадывающему. Затем предложим ему отгадать первую букву

отрывка. Если догадка оказалась правильной, то об этом сообщается отгадывающему и предлагается определить вторую букву. Если первая буква не отгадана правильно, то она также сообщается, и переходят к следующему отгадыванию. Это продолжается до конца текста. По ходу эксперимента отгадывающий выписывает правильный текст вплоть до последней буквы для использования его при отгадывании очередной буквы. Результат эксперимента такого рода приводится ниже. Пробел между словами считается дополнительной буквой, т. е. имеется 27-буквенный алфавит. В строках, помеченных (1), написан исходный текст. В строках, помеченных (2), на месте угаданных букв проставляется черта, а в случае неправильного отгадывания выписывается буква исходного текста.

(1) THE ROOM WAS NOT VERY LIGHT A SMALL OBLONG

(2) ---- ROO-----NOT-V-----I-----SM----OBL-----

(1) READING LAMP ON THE DESK SHED CLOW ON

(2) REA-----O-----D----SHED-CLO--O-

(1) POLISHED WOOD BUT LESS ON THE SHABBY RED CARPET

(2) P-L-S-----O---BU--L-S--O-----SH-----RE-C----- (8)

Из общего числа 129 букв 89 букв, или 69%, были отгаданы правильно. Ошибки, как и следовало ожидать, встречались наиболее часто в начале слов и слогов, где ход мысли может быть наиболее разнообразным. Можно было бы подумать, что вторые строки в (8), которые назовем приведенным текстом, содержат гораздо меньше информации, чем первые. В действительности же обе строки содержат одинаковую информацию в том смысле, что возможно, по крайней мере в принципе, определить первую строку по второй. Для выполнения этого тому, кто отгадывает текст, необходим двойник. Двойник (который если не биологически, то по крайней мере математически идентичен с оригиналом) должен будет отвечать точно так же, как отгадывающий, когда столкнется с аналогичной ситуацией. Предположим теперь, что имеется только приведенный текст (8). Двойнику предлагается отгадать исходный текст. В каждый момент нам известно, правильно выполнено отгадывание или нет, поскольку он отгадывает так же хорошо, как и его предшественник, и присутствие черточки в приведенном тексте соответствует правильному отгадыванию. Буквы, которые он отгадывает неправильно, также доступны, так что на каждом шагу ему может быть предоставлена в точности та информация, которая была получена первым отгадывающим.

От необходимости в двойнике в таком мыслимом эксперименте можно отказаться следующим приемом. Вообще говоря, хорошее предсказание не требует знания более чем N букв предшествующего текста, где N не очень велико. Имеется конечное число воз-

работе с нормативным литературным английским текстом. Статьи из газеты, научные работы или поэзия требуют меньшего числа попыток.

Приведенный текст в этом случае также содержит ту же информацию, что и исходный. Опять, используя двойника, предлагаем ему на каждой стадии отгадывать текст столько раз, сколько единиц соответствует цифре в приведенном тексте, и таким путем обнаружить исходный текст. Для того чтобы исключить субъективный элемент, можно просить отгадывающего указать для каждой N -граммы текста наиболее вероятную букву, вторую по вероятности букву и т. д. Эта совокупность данных может служить как для предсказания, так и для отгадывания.

Как и раньше, приведенный текст можно рассматривать как закодированный вариант исходного. Именно английский язык с алфавитом в 27 символов, А, В, ..., Z, пробел, переведен на новый язык с алфавитом 1, 2, ..., 27. Перевод выполнен таким образом, что символ 1 имеет теперь наибольшую частоту. Символы 2, 3, 4 имеют последовательно все меньшие и меньшие частоты, и заключительные символы 20, 21, ..., 27 встречаются вообще очень редко. Таким образом, перевод в значительной степени упростил рассматриваемую статистическую структуру. Избыточность, которая проявлялась вначале в сложных зависимостях между группами букв, теперь свелась к значительной разнице между вероятностями новых символов. Это, как будет видно позднее, явится основой для оценки энтропии с помощью таких экспериментов.

Для определения того, как возможность предсказания зависит от числа n предшествующих букв, известных отгадывающему, был выполнен более сложный эксперимент. Сотня выборок по 15 букв английского текста была выбрана случайным образом из книги. Отгадывающему было предложено отгадывать текст по одной букве в каждой выборке, как в предыдущем эксперименте. Таким образом, была получена сотня выборок, где отгадывающий знал 0, 1, 2, 3, ..., 14 предшествующих букв. Отгадывающему могут быть предоставлены все средства, которые он пожелает, например, различные статистические таблицы, однобуквенные, двухбуквенные и трехбуквенные таблицы, таблицы частот начальных букв в словах, сводка частот наиболее употребительных слов и любой словарь. Выборки для эксперимента были взяты из книги «Виргинец Джефферсон» Дюма Малона. Полученные результаты вместе с аналогичным экспериментом, когда отгадывающему были известны 100 букв, собраны в табл. 1. (см. стр. 678—679). Столбец соответствует числу известных предшествующих букв; номер строки указывает число отгадываний. На пересечении N -го столбца и S -й строки стоит число раз, при которых опознавание правильной буквы произошло при S -м отгадывании, когда известны предыдущие $N - 1$ букв.

Например, цифра 19 на пересечении 6-го столбца и 2-й строки означает, что при известных пяти предшествующих буквах правильная буква была получена в девятнадцати случаях из ста при втором угадывании. Первые два столбца этой таблицы были получены не экспериментально, как описано выше, а были вычислены непосредственно с помощью известных частот отдельных букв и диграмм. Таким образом, если ни одна из предшествующих букв не задана, то наиболее вероятным символом является промежуток между словами (вероятность 0,182) в случае, если первое угадывание оказалось ошибочным, следующим должно быть Е (вероятность 0,107) и т. д. Эти вероятности суть частоты, с которыми правильные угадывания произойдут при первом, втором и т. д. эксперименте при наилучшем предсказании. Подобным же образом простое вычисление при помощи таблицы диграмм дает результаты для 1-го столбца. Поскольку частотные таблицы определены по весьма длинным выборкам из английского текста, то эти два столбца подчинены меньшим статистическим ошибкам, чем остальные.

Далее будет показано, что предсказание действительно улучшается, если не учитывать статистических флуктуаций, зависящих от знания прошлого; так, это подтверждается большим количеством правильных угадываний и меньшим числом угадываний, требующих много проб.

Был проделан также один эксперимент с «обратным» предсказанием, когда приходилось угадывать букву, предшествующую уже известным. Хотя задача субъективно гораздо более трудная, но результаты были ненамного хуже. Так, для выборки в 101 букву из той же самой книги были получены следующие результаты:

Номер угадывания	1	2	3	4	5	6	7	8	8
Вперед	70	10	7	2	2	3	3	0	4
Назад	66	7	4	4	4	2	1	2	9

Следует учесть, что N -граммная энтропия F_N для обратного языка равна N -граммной энтропии для прямого языка, что можно видеть из второго выражения в уравнении (1). Оба члена имеют одно и то же значение в прямом и обратном случаях.

4. Идеальное N -граммное предсказывание

Данные табл. 1 могут быть использованы для получения верхней и нижней границ для N -граммной энтропии F_N . Для этого необходимо вначале вывести некоторые общие закономерности, относящиеся к наилучшему возможному предсказанию языка, когда предыдущие N букв известны. Для языка имеются условные вероятности $p_{i_1, i_2, \dots, i_{N-1}}(j)$, обозначающие вероятность того, что после $(N-1)$ -граммы $i_1, i_2, \dots, i_{N-1}$ будет стоять буква j .

Таблица 1

	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	100
1	18,2	29,2	36	47	51	58	48	66	66	67	62	58	66	72	60	80
2	10,7	14,8	20	18	13	19	17	15	13	10	9	14	9	6	18	7
3	8,6	10,0	12	14	8	5	3	5	9	4	7	7	4	9	5	—
4	6,7	8,6	7	3	4	1	4	4	4	4	5	6	4	3	5	3
5	6,5	7,1	1	1	3	4	3	6	1	6	5	2	3	—	—	4
6	5,8	5,5	4	5	2	3	2	—	—	1	4	2	3	4	1	2
7	5,6	4,5	3	3	2	2	8	—	1	1	1	4	1	—	4	1
8	5,2	3,6	2	2	1	1	2	1	1	1	1	—	2	1	3	—
9	5	3	4	—	5	1	4	—	2	1	1	2	—	1	—	1
10	4,3	2,6	2	1	3	—	3	1	—	—	—	—	2	—	—	—
11	3,1	2,2	2	2	2	1	—	—	1	3	—	1	1	2	1	—
12	2,8	1,9	4	—	2	1	1	1	—	—	2	1	1	—	1	1
13	2,4	1,5	1	1	1	1	1	1	1	1	—	1	1	—	—	—

[illegible]

Наилучшее угадывание за уже известной $(N-1)$ -граммой будет для буквы с наибольшей условной вероятностью. Второй должна быть угадана буква со следующей по величине вероятностью и т. д. Машина или лицо, производящее угадывание наилучшим способом, будет угадывать буквы в порядке убывания условных вероятностей. Итак, процесс приведения текста с таким идеальным предсказанием состоит в отображении букв на числа от 1 до 27 таким образом, что наиболее вероятная следующая буква (при условии, что известна предыдущая $N-1$ -грамма) отображается в 1 и т. д. Частота единиц в приведенном тексте будет даваться выражением

$$q_1^N = \sum p(i_1, i_2, \dots, i_{N-1}, j), \quad (10)$$

где сумма берется по всем $(N-1)$ -граммам, $i_1, \dots, i_{N-1}$ и j — буква, максимизирующая p для данной $(N-1)$ -граммы. Аналогично частота цифры 2, обозначаемая q_2^N , дается той же самой формулой, но j выбирается так, что она дает второе по величине значение p и т. д.

При таком отображении N -грамм получается другая совокупность вероятностей для символов приведенного текста: $q_1^{N+1}, q_2^{N+1}, \dots, q_{27}^{N+1}$. Поскольку предсказание основано теперь на большем знании прошлого, то следует ожидать, что вероятности маловероятных букв будут больше; и действительно, можно доказать следующее неравенство:

$$\sum_{i=1}^s q_i^{N+1} \geq \sum_{i=1}^s q_i^N, \quad s = 1, 2, \dots \quad (11)$$

Оно означает, что вероятность правильного угадывания за первые s попыток, когда известны предыдущие N букв, больше или равна той же вероятности, когда известны предыдущие $N-1$ букв для всех s . Для доказательства представим себе, что вероятности $p(i_1, i_2, \dots, i_N, j)$ упорядочены в таблицу, где номера j расположены по горизонталям, а все N -граммы — по вертикали. Таблица будет содержать, таким образом, 27 столбцов и 27^N строк. Член в левой части неравенства (11) есть сложенная по всем строкам сумма s наибольших чисел в каждой строке. Член, стоящий в правой части неравенства (11), также есть сумма чисел таблицы, причем из каждой строки берется также s чисел, но не обязательно наибольших. Это следует из того, что член, стоящий справа в выражении (11), может быть вычислен из аналогичной таблицы для $(N-1)$ -грамм, полученной из таблицы N -грамм. Каждая строка в $(N-1)$ -граммной таблице есть сумма 27 строк в N -граммной таблице, поскольку

$$p(i_2, i_3, \dots, i_N, j) = \sum_{i_1=1}^{27} p(i_1, i_2, \dots, i_N, j). \quad (12)$$

Сумма s наибольших чисел некоторой строки $(N-1)$ -граммной таблицы будет равна сумме $27s$ чисел, отобранных из соответствующих 27 строк N -граммной таблицы только в том случае, когда эти отобранные строки попадают лишь в s столбцов таблицы. Для того чтобы равенство в выражении (15) было верно для заданного s , это утверждение должно выполняться для каждой строки $(N-1)$ -граммной таблицы. В этом случае первая буква N -граммы не влияет на множество s наиболее вероятных выборов для следующей буквы, хотя упорядочение внутри этого множества может от нее зависеть. Однако если равенство в выражении (11) выполняется для всех s , то упорядочение также будет независимым от первой буквы N -граммы. Приведенный текст, полученный с помощью $(N-1)$ -граммного идеального предсказания, совпадает тогда с текстом, полученным от идеального N -граммного предсказания.

Поскольку частные суммы

$$O_s^N = \sum_{i=1}^n q_i^N, \quad s = 1, 2, \dots \quad (13)$$

монотонно не убывают в зависимости от N и меньше единицы для всех N , то они должны стремиться к пределу при $N \rightarrow \infty$, т. е. q_i^N стремится к пределу q_i^∞ . Эти пределы могут быть интерпретированы как относительная частота первого, второго и т. д. угадывания при известной всей (бесконечной) предыстории текста.

Идеальное N -граммное предсказание можно рассматривать, как уже было указано, как преобразователь, переводящий язык в последовательность чисел от 1 до 27. Он обладает при этом следующими двумя свойствами.

1. Выходной символ есть функция от входного символа в настоящий момент времени и от предшествующих $(N-1)$ букв.

2. Он мгновенно обратим. Входной сигнал может быть восстановлен без потери времени с помощью соответствующей операции над приведенным текстом. В самом деле обратный оператор также действует только над $(N-1)$ предшествующими символами приведенного текста и выходным символом, полученным в данный момент.

Все изложенное является доказательством того, что частоты выходных символов при $(N-1)$ -граммном предсказании удовлетворяют неравенствам

$$\sum_{i=1}^s q_i^N \geq \sum_{i=1}^s q_i^{N-1}, \quad s = 1, 2, \dots, 27 \quad (14)$$

и могут быть приложимы к любому преобразователю с перечисленными двумя свойствами. В самом деле можно снова представить себе таблицу с различными $(N-1)$ -граммами, расположенными по вертикали, и буквой, появляющейся в данный момент на этой вертикали. Так как выходной символ в данный момент зависит только

от этих величин, то получаем определенный выходной символ, который можно расположить на пересечении соответствующих строки и столбца. Далее, мгновенная обратимость требует, чтобы на каждой строке не было одинаковых чисел. Иначе при обратном декодировании возникала бы неопределенность в выборе между несколькими возможностями для входной буквы, принимаемой в настоящий момент. Общая вероятность s наиболее вероятных символов на выходе, скажем $\sum_1^s r_i$, будет суммой чисел, просуммированной по всем строкам, взятым по s из каждой строки, и, следовательно, не больше чем сумма s наибольших чисел в каждой строке. Таким образом, имеем

$$\sum_1^s q_i^N \geq \sum_1^s r_i, \quad s = 1, 2, \dots, 27. \quad (15)$$

Иными словами, идеальное предсказывающее устройство, которое было описано выше, обладает преимуществом перед всеми преобразователями, которые могут быть применены к языку, и обладают двумя указанными свойствами. Грубо говоря, идеальное предсказывающее устройство разбивает вероятности различных символов на меньшие группы, чем любой другой мгновенно обратимый преобразователь, пользующийся тем же самым числом букв.

Совокупности чисел, удовлетворяющих неравенству (15), изучались Мюирхедом в связи с теорией алгебраических неравенств¹⁾. Если неравенство (15) выполняется, когда q_i^N и r_i расставлены в порядке убывания их величин, и $\sum_1^{27} q_i^N = \sum_1^{27} r_i^N$ (в данном случае это выполняется, поскольку общая вероятность в каждом случае равна 1), то первое множество чисел q_i^N мажорирует второе множество чисел r_i . Известно, что свойство мажорирования эквивалентно каждому из следующих свойств.

1. Числа r_i могут быть получены из q_i^N с помощью конечного числа «переносов». Под переносом понимается передача вероятности от больших q к меньшим подобно тому, как тепло течет от более нагретых к менее нагретым телам, а не наоборот.

2. Числа r_i могут быть получены из q_i^N обобщенным «усреднением», т. е. существуют неотрицательные действительные числа a_{ij} , для которых $\sum_j a_{ij} = \sum_i a_{ij} = 1$, такие, что

$$r_i = \sum_j a_{ij} (q_j^N). \quad (16)$$

¹⁾ Hardy G., Littlewood J., Polya G., *Inequalities*, Cambridge University Press, 1934. (Русский перевод: Харди Г. Г., Литтлвуд Дж. Е., Полиа Г., *Неравенства*, ИЛ, 1948.—Прим. ред.).

5. Граница для энтропии на основании частот предсказания

Если известны частоты символов в приведенном тексте, полученном с помощью N -граммного идеального предсказания, то можно получить верхние и нижние оценки для N -граммной энтропии F_N исходного текста. Эти границы следующие:

$$\sum_{i=1}^{27} i (q_i^N - q_{i+1}^N) \log i \leq F_N \leq - \sum_{i=1}^{27} q_i^N \log q_i^N. \quad (17)$$

Верхняя оценка следует немедленно из того факта, что максимальная возможная энтропия в языке с частотами букв q_i^N равна $-\sum q_i^N \log q_i^N$. Таким образом, энтропия на символ приведенного текста не больше этой величины. N -граммная энтропия приведенного текста равна N -граммной энтропии исходного языка, что легко усмотреть из выражения (1) для F_N . Имеющиеся там суммы будут содержать в точности те же члены, хотя, может быть, в другом порядке. Эта верхняя оценка справедлива независимо от того, будет ли предсказание идеальным.

Получить нижнюю оценку более трудно. Необходимо показать, что при любом выборе N -граммных вероятностей $p(i_1, \dots, i_N)$ справедливо неравенство

$$\sum_{i=1}^{27} i (q_i^N - q_{i+1}^N) \log i \leq \sum_{i_1, \dots, i_N} p(i_1, \dots, i_N) \log p_{i_1 \dots i_{N-1}}(i_N). \quad (18)$$

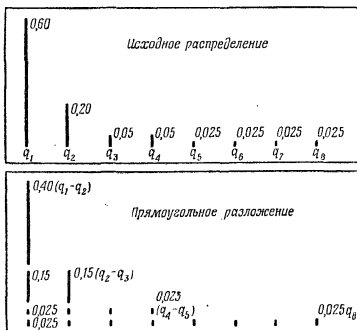
Слагаемые, стоящие в левой части неравенства, могут быть интерпретированы следующим образом: представим себе числа q_i^N в виде упорядоченной последовательности отрезков убывающей высоты (рис. 3). Истинное значение q_i^N можно рассматривать как сумму прямоугольных распределений, как это показано на том же рисунке. Левая часть неравенства (18) есть тогда энтропия этого множества отрезков. Таким образом, i -е прямоугольное распределение имеет общую вероятность $i(q_i^N - q_{i+1}^N)$. Энтропия этого распределения равна $\log i$. Общая энтропия равна

$$\sum i (q_i^N - q_{i+1}^N) \log i. \quad (19)$$

Задача, таким образом, свелась к тому, чтобы показать, что система вероятностей $p(i_1, \dots, i_N)$ с наилучшими частотами предсказания q_i имеет энтропию F_N , большую или равную энтропии прямоугольной системы, полученной из того же самого множества q_i .

Как указывалось выше, q_i получены из $p(i_1, \dots, i_N)$ упорядочиванием каждой строки таблицы в убывающем порядке по величине и сложением по вертикали. Таким образом, каждое q_i можно пред-

ставить в виде суммы совокупности монотонно убывающих распределений. Каждое из этих распределений заменим его прямоугольным разложением. Каждое из них заменится, вообще говоря, 27 такими распределениями. При этом каждое q_i окажется суммой



Р и с. 3. Прямоугольное разложение монотонного распределения.

27×27^N распределений, содержащих от одного до 27 элементов, расположенных, начиная с левого столбца. Энтропия этого множества меньше или равна энтропии исходного множества распределений, поскольку почленное сложение двух или более распределений всегда увеличивает энтропию. Это есть конкретное применение общей теоремы, состоящей в том, что $H_y(x) \leq H(x)$, и выполненной для любых случайных величин x и y . Равенство выполняется тогда и только тогда, когда складываемые распределения пропорциональны. Теперь можно складывать различные компоненты одной длины без изменения энтропии (поскольку в этом случае распределения пропорциональны). В результате приходим, исходя из начальных N -граммных вероятностей, к прямоугольному разложению для q_i с помощью ряда процессов, которые уменьшают или сохраняют энтропию. Следовательно, энтропия исходной системы F_N больше или равна энтропии прямоугольного разложения q_i . Этим заканчивается доказательство.

Следует отметить, что нижняя граница строго меньше F_N , за исключением того случая, когда каждая строка таблицы имеет

прямоугольное распределение. Это возможно тогда и только тогда, когда для каждой возможной $(N-1)$ -граммы имеется ряд равновероятных возможностей для следующей буквы, а остальные возможности имеют вероятности, равные нулю.

Покажем, что верхние и нижние границы для F_N , даваемые формулой (17), являются монотонно убывающими функциями от N . Это очевидно для верхней оценки, поскольку q_i^{N+1} мажорируют q_i^N , и любой выравнивающий перенос в множестве вероятностей увеличивает энтропию. Для доказательства того, что нижняя оценка также монотонно убывает, покажем, что величина

$$u = \sum i(q_i - q_{i+1}) \log i \quad (20)$$

увеличивается при выравнивающем переносе между q_i . Предположим, что перенос происходит от q_i к q_{i+1} , причем первая вероятность убывает на Δq , а вторая возрастает на ту же величину. Тогда в сумме (20) изменяются три слагаемых и изменение u дается выражением

$$\Delta u = [- (i-1) \log (i-1) + 2i \log i - (i+1) \log (i+1)] \Delta q. \quad (21)$$

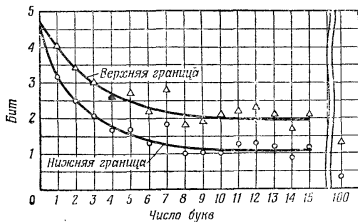
Член в скобках имеет вид $-f(x-1) + 2f(x) - f(x+1)$, где $f(x) = x \log x$. Но функция $f(x)$ выпукла при положительных x , поскольку $f''(x) = \frac{1}{x} > 0$. Член в скобках является второй разностью между ординатой кривой при $x=i$ и ординатой средней точки хорды, соединяющей $i-1$ и $i+1$, и, следовательно, отрицателен. Поскольку Δq также отрицательно, то изменение u под действием переноса отрицательно. Даже более простые рассуждения показывают, что это также верно для переноса от q_1 к q_2 или от q_{26} к q_{27} (когда меняются только два слагаемых). Тем самым получаем, что нижняя оценка, основанная на N -граммных частотах предсказания q_i^N , больше или равна оценке, вычисленной по $(N+1)$ -граммным частотам предсказания q_i^{N+1} .

6. Результаты эксперимента для английского языка

На основании данных табл. 1 были вычислены верхние и нижние границы в неравенстве (17). Данные были сначала несколько сглажены для уничтожения выборочных отклонений. Малым числам в этой таблице можно доверять меньше всего, и они были усреднены по группам. Таким образом, в 4-м столбце числа 47, 18 и 14 не были изменены, а для оставшихся строк от (4-й до 20-й) сумма, равная 21, была разделена по ним равномерно. Верхние границы, даваемые неравенством (17), были затем вычислены для каждого столбца и получены результаты, приведенные ниже.

столбец	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	100
верхняя оценка	4,03	3,42	3	2,6	2,7	2,2	2,8	1,8	1,9	2,1	2,2	3,2	1	1,7	2,1	1,3
нижняя оценка	3,19	2,5	2,1	1,7	1,7	1,3	1,8	1	1	1	1,3	1,3	1,2	0,9	1,2	0,6

Очевидно, что в этих числах имеется еще значительная статистическая ошибка, связанная с идентификацией наблюдаемых выборочных частот с вероятностями предсказания. Надо также вспомнить, что нижняя граница была выведена только для устройства идеального предсказания, в то время как использованные нами частоты получены из предсказания, сделанного человеком. Некоторые грубые вычисления показывают, однако, что расхождение



Р и с. 4. Верхняя и нижняя границы, полученные экспериментальным путем для энтропии английского языка с 27-буквенным алфавитом.

между действительным значением F_N и нижней оценкой, полученной идеальным устройством предсказания (объясняющаяся тем, что условные вероятности не имели прямоугольного распределения), вполне компенсирует неумение человека предсказывать идеально.

Таким образом, можно в достаточной степени уверенно сказать, что обе группы независимы от статистических ошибок.

Полученные значения нанесены под рубрикой N на рис. 4.

Автор благодарен миссис Мэри Шэннон и д-ру Оливеру за помощь в экспериментальной работе и за большое число советов и критических замечаний по поводу теоретической части статьи.

УПРОЩЕННЫЙ ВЫВОД ЛИНЕЙНОЙ ТЕОРИИ СГЛАЖИВАНИЯ И ПРЕДСКАЗАНИЯ ПО МЕТОДУ НАИМЕНЬШИХ КВАДРАТОВ¹⁾

Краткое содержание

Основные результаты теории сглаживания и предсказания стационарных временных рядов Винера — Колмогорова получены новым методом. Примененный подход навеян физическими соображениями, основанными на теории электрических цепей, и не требует применения интегральных уравнений или функций корреляции. Рассмотрены случаи сглаживания с бесконечной задержкой, случай чистого предсказания (без шума) и общая проблема сглаживания и предсказания. В конце обсуждаются основные предположения теории для выяснения вопроса об условиях ее адекватности и для того, чтобы предупредить ее необоснованные приложения.

1. Введение

В классическом отчете, написанном для Национального совета оборонных исследований (США), Винер²⁾ изложил математическую теорию сглаживания и предсказания, представляющую значительный интерес для теории связи. Почти в то же самое время эта теория была независимо развита Колмогоровым³⁾. К сожалению, работы Винера и Колмогорова содержат несколько устрашающий математический аппарат — винеровский отчет в желтой обложке приобрел среди смущенных инженеров название «желтой опасности» — и это воспрепятствовало широкому распространению и использованию, которых заслуживает вышеупомянутая теория. В дан-

¹⁾ Bode H., Shannon C., A simplified derivation of linear least square smoothing and prediction theory, *Proc. IRE*, 38 4, (1950), 417.

²⁾ Wiener N., The interpolation, extrapolation and smoothing of stationary time series, Wiley, New York, 1949.

³⁾ Колмогоров А. Н., Интерполирование и экстраполирование стационарных случайных последовательностей, *Изв. АН СССР (сер. матем.)*, 5 (1941), 3. (Впервые эти результаты были опубликованы Колмогоровым в 1939 г., т. е. за несколько лет до начала работ Н. Винера по этой тематике; см. Колмогоров А. Н., Sur l'interpolation et extrapolation de suites stationnaires, *C. R. Acad. Sci.*, 208 (1939), 2043—2045. — *Прим. ред.*)

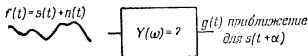
ной статье будут получены основные результаты теории сглаживания новым методом, хотя и не таким строгим и общим, как методы Винера и Колмогорова, но наиболее простым, в особенности для читателей, работающих в области электрических цепей. Математические выводы в настоящей трактовке большей частью имеют прямое физическое истолкование, которое позволяет интуитивно осмыслить ход математических рассуждений.

2. Проблема и основные предположения

Основная проблема, которая будет рассмотрена, может быть сформулирована следующим образом: дан искаженный сигнал $f(t)$, который является суммой истинного сигнала $s(t)$ и искажающего шума $n(t)$:

$$f(t) = s(t) + n(t).$$

Желательно так обработать $f(t)$, чтобы получить по мере возможности истинный сигнал $s(t)$. В более общем случае желательно объединить эту операцию сглаживания с предсказанием, т. е.



Р и с. 1. Проблема сглаживания и предсказания.

обработать $f(t)$ таким образом, чтобы получить хорошее приближение к тому, что $s(t)$ будет представлять собой в будущем, скажем через α секунд, или каким он был в прошлом, α секунд тому назад. В этом случае желательно аппроксимировать $s(t + \alpha)$, где α соответственно положительно или отрицательно. Общая картина схематически изображена на рис. 1; задача заключается в нахождении звена, обозначенного «?».

Мы увидим, что эта проблема и ее обобщения имеют широкое применение не только в теории связи, но и в таких разнообразных областях, как экономическое предсказание, предсказание погоды, артиллерия, статистика и т. д.

Теория Винера — Колмогорова базируется на трех основных предположениях, которые определяют область применения результатов этой теории. Эти предположения следующие.

1. Временные последовательности, представляющие сигнал $s(t)$ и шум $n(t)$, *стационарны*. Это значит по существу, что статические свойства сигнала и шума не изменяются со временем. Теория не может быть применена, например, к *долговременным* экономи-

ческим явлениям, так как статистика, скажем рыночных цен сегодня, не та же, что в 1850 г.

2. В качестве критерия ошибки приближения взято *средне-квадратичное различие* между действительным и желаемым откликом искомого звена. Согласно рис. 1, это означает, что надо найти звено «?» так, чтобы минимизировать средне-квадратичную ошибку $[g(t) - s(t + \alpha)]^2$, среднее берется по всем возможным функциям сигнала и шума, взвешенным соответственно вероятностям их появления. Это называется средним по ансамблю.

3. Операция для предсказания и сглаживания предполагается *линейной* операцией над имеющейся информацией, или — в терминах теории связи — искомое звено должно быть линейным физически осуществимым фильтром. Имеющаяся в распоряжении информация состоит в знании прошлого искаженного сигнала, т. е. функции $f(t)$ при $t \leq t_1$, где t_1 — настоящий момент времени. Линейный физически реализуемый фильтр производит линейную операцию над $f(t)$ во всей этой области, как увидим далее, в соответствии с равенствами (3) и (4).

Поэтому теория может быть охарактеризована как *линейное предсказание и сглаживание по методу наименьших квадратов стационарных временных рядов*. Ясно, что теория применима только тогда, когда эти три предположения удовлетворяются или хотя бы приблизительно удовлетворяются. Если одно из них изменено или устранено, проблема предсказания и сглаживания становится математически очень трудной, и нам мало что известно о точных решениях. Некоторые из ограничений, налагаемых сделанными предположениями, будут обсуждены далее.

Как можно предсказать все будущее поведение функции, когда все, что известно, — это только искаженные данные о ее прошлом? Этот вопрос тесно связан с проблемами причинности и индукции в философии и вопросом о смысле физических законов. В общем физическое предсказание зависит в основном от *предположения*, что закономерности, наблюдавшиеся в прошлом, будут сохранены в будущем. Это предположение не может быть доказано дедуктивно, т. е. чисто математической аргументацией, так как можно легко представить себе математический мир, в котором это предположение не верно. Так же это не может быть установлено индуктивно, т. е. обобщением из эксперимента, так как такое обобщение будет основано на предположениях, которые мы попытаемся установить. Это предположение может рассматриваться только как один из центральных постулатов физики.

Классическая физика стремилась свести физический мир к системе строгих законов причинности. Будущее физической системы при этом точно предсказуемо из знания его прошлого, и все, что требуется — это знание настоящего положения системы. Сове-

менная квантовая физика заставила нас отбросить эти представления как несостоятельные. Законы физики теперь истолковываются как статистические, и единственно возможное предсказание является статистическим. «Точные» законы физики содержат неопределенности; эти неопределенности малы, когда изучаются большие объекты, и относительно велики для объектов атомного масштаба.

Теория линейного минимально-квадратичного предсказания и сглаживания основывается на статистическом предсказании. Основное предположение, что статистические закономерности в прошлом должны сохраниться в будущем, проявляется в математике как предположение, что сигнал и шум являются *стационарными* временными рядами. Это означает, например, что статистический параметр сигнала, усредненный по всему прошлому, даст ту же самую величину, как этот же параметр, усредненный по будущему.

Предсказание существенно зависит от наличия корреляции между будущей величиной сигнала $s(t_1 + \alpha)$, где t_1 — настоящий момент времени, и известной величиной $f(t) = s(t) + n(t)$ для $t \leq t_1$. Предположение, что предсказание является *линейной* операцией, влечет за собой использование только одного типа корреляции, а именно линейной корреляции, т. е. $\overline{s(t_1 + \alpha) f(t)}$. Если бы эта корреляция была нулевой для всех $t \leq t_1$, то, как будет показано ниже, никакое существенное линейное предсказание невозможно. Тогда наилучшая средне-квадратичная оценка для $s(t_1 + \alpha)$ будет нулем.

3. Свойства линейных фильтров

Для удобства в этом разделе будут резюмированы некоторые хорошо известные результаты, относящиеся к фильтрам. Линейный фильтр может характеризоваться двумя различными, но эквивалентными способами. Первое и более распространенное — описание в терминах комплексного коэффициента передачи $Y(\omega)$. Если чистая синусоида частоты ω_1 и амплитуды E воздействует на вход фильтра, откликом является также синусоида частоты ω_1 и амплитуды $|Y(\omega_1)|E$. Фаза на выходе определяется углом $Y(\omega_1)$ — фазой фильтра на этой частоте. Часто удобно записывать комплексный коэффициент передачи $Y(\omega)$ в форме $Y(\omega) = e^{A(\omega)} e^{iB(\omega)}$, где $A(\omega) = \log |Y(\omega)|$ — усиление, а $B(\omega)$ — угол $[Y(\omega)]$ является фазой. Предположим, что фильтр может содержать идеальный усилитель так же, как и пассивные элементы; поэтому можно прибавить любую постоянную к A , чтобы сделать абсолютный уровень таким, как нужно.

Второй способ описания свойств фильтра вводится посредством функций времени. Пусть $K(t)$ — обратное преобразование

Фурье для $Y(\omega)$:

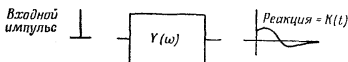
$$K(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} Y(\omega) e^{i\omega t} d\omega. \quad (1)$$

Тогда $Y(\omega)$ — прямое преобразование Фурье от $K(t)$

$$Y(\omega) = \int_{-\infty}^{\infty} K(t) e^{-i\omega t} dt. \quad (2)$$

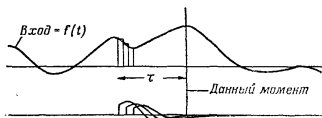
Знание $K(t)$ полностью эквивалентно знанию $Y(\omega)$, одно из них может быть вычислено, если известно второе.

Временная функция $K(t)$ представляет реакцию фильтра на единичный импульс, приложенный к входу в момент $t=0$, как



Р и с. 2. Импульсная реакция фильтра.

показано на рис. 2. Из этого можно легко получить реакцию фильтра на произвольный входной сигнал $f(t)$. Нужно просто разделить входной сигнал на большое число вертикальных ступенек, как показано на рис. 3. Каждая ступенька может рассматриваться как импульс величины $f(t) dt$, который производит



Р и с. 3. Реакция на произвольное воздействие как сумма импульсных реакций.

реакцию $f(t) dt K(t_1 - t)$ в любой последующий момент t_1 . После суммирования вклада всех ступеней получаем хорошо известную формулу

$$g(t_1) = \int_{-\infty}^{t_1} f(t) K(t_1 - t) dt \quad (3)$$

для общей реакции в момент t_1 .

Для изучения теории сглаживания выражение (3) может быть заменено несколько иной формулой. Подставляя $\tau = t_1 - t$, мы имеем

$$g(t_1) = \int_0^{\infty} f(t_1 - \tau) K(\tau) d\tau. \quad (4)$$

В этой формуле τ означает время запаздывания, так что $f(t_1 - \tau)$ дает значение входного сигнала τ секунд назад. $K(\tau)$ является функцией, подобной импульсной проводимости, однако отнесенной в прошлое, а не в будущее, как показано на рис. 4. Очевидно, это

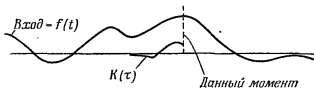


Рис. 4. Реакция фильтра как взвешенное среднее по прошедшим значениям входного сигнала.

есть весовая функция, на которую должно быть умножено входное напряжение, чтобы определить его вклад в реакцию в данный момент.

Критерий физической осуществимости будет приведен либо в терминах K , либо в терминах функции Y . В терминах импульсной реакции $K(t)$ необходимо, чтобы $K(t)$ имело нулевое значение для $t < 0$, т. е. фильтр не может реагировать на импульс до того, как импульс возник. Кроме того, $K(t)$ должно стремиться к нулю (с разумной скоростью) при $t \rightarrow +\infty$. Действие импульса, очевидно, должно затухать со временем.

Эти же требования сохраняют смысл при трактовке K как весовой функции. Фильтр не может применять взвешивания к той части входного сигнала, которая еще только должна появиться, поэтому $K(\tau) = 0$ при $\tau < 0$. Действие очень удаленных прошлых значений будет постепенно затухать, так что $K(\tau)$ должно стремиться к нулю, если $\tau \rightarrow \infty$. Следует отметить, что эти условия являются также достаточными для физической осуществимости фильтра в том смысле, что можно получить сколь угодно близкое приближение к данной импульсной реакции $K(t)$ при помощи пассивной цепи с сосредоточенными постоянными с одним усилителем.

В терминах частотной характеристики главным условием физической осуществимости является то, что $Y(\omega)$, рассматриваемая как функция комплексной переменной ω , должна быть аналитической функцией в полуплоскости, определяемой условием $\text{Im}(\omega) < 0$.

Кроме того, функция должна вести себя на действительной оси частот¹⁾ так, чтобы интеграл (5) был конечен.

$$\int_0^{\infty} \frac{\log |Y(\omega)|}{1+\omega^2} d\omega. \quad (5)$$

Требования физической осуществимости приводит к хорошо известным соотношениям между амплитудно-частотной и фазово-частотной характеристиками. Для данного усиления $A = \log |Y\omega|$, удовлетворяющего (5), существует минимально-фазовая характеристика. Эта фаза определяется так:

$$B(\omega_0) = \frac{2\omega_0}{\pi} \int_0^{\infty} \frac{A(\omega) - A(\omega_0)}{\omega^2 - \omega_0^2} d\omega. \quad (6)$$

Если квадрат требуемого усиления $|Y(\omega)|^2 = Y(\omega) \bar{Y}(\omega)$ — рациональная функция частоты ω , скажем $P_1(\omega)/P_2(\omega)$, где $P_1(\omega)$ и $P_2(\omega)$ — полиномы, то минимально-фазовая характеристика может быть получена следующим образом. Вычислим корни $P_1(\omega)$ и $P_2(\omega)$ и запишем $|Y(\omega)|^2$ как

$$|Y(\omega)|^2 = k^2 \frac{(\omega - \alpha_1)(\omega - \bar{\alpha}_1)(\omega - \alpha_2)(\omega - \bar{\alpha}_2) \dots}{(\omega - \beta_1)(\omega - \bar{\beta}_1)(\omega - \beta_2)(\omega - \bar{\beta}_2) \dots}, \quad (7)$$

где все $\alpha_1, \alpha_2, \dots, \beta_1, \beta_2, \dots$ имеют мнимые части больше нуля, т. е. имеются нули и полюсы $|Y(\omega)|^2$ в верхней полуплоскости, а сопряженные члены имеют соответствующие корни и полюсы в нижней полуплоскости. Фильтр с минимально-фазовой характеристикой имеет тогда коэффициент передачи, равный

$$Y(\omega) = k \frac{(\omega - \alpha_1)(\omega - \alpha_2) \dots}{(\omega - \beta_1)(\omega - \beta_2) \dots}. \quad (8)$$

Минимально-фазовый фильтр имеет то важное свойство, что его обращение с коэффициентом передачи $Y^{-1}(\omega)$ также физически осуществимо²⁾. Если пропустить сигнал $f(t)$ через фильтр $Y(\omega)$, можно восстановить его в первоначальной форме, пропустив этот сигнал через обратный фильтр. Более того, восстановление происходит *без потери времени*. С другой стороны, не существует физи-

¹⁾ Включая точку на бесконечности. Реальные физические фильтры всегда имеют нулевое усиление на бесконечной частоте, а вышеприведенное требование показывает, что стремление к нулю не может быть слишком быстрым. Стремление типа ω^{-n} (6 дБ на октаву) допустимо, но $e^{-|\omega|}$ или $e^{-\omega^2}$ делает интеграл (5) расходящимся, что приводит к физической неосуществимости.

²⁾ Если первоначальная функция имеет нуль на бесконечности, так что обращение имеет там полюс, возникают осложнения, но в реальных случаях может быть получено подходящее приближение.

чески осуществимого точного обращения для неминимально-фазового фильтра. Самое лучшее, что можно сделать, — создать устройство, которое обладало бы всеми свойствами теоретического обращения, если не учитывать дополнительного фазового запаздывания. Дополнительное фазовое запаздывание может быть сбалансировано для получения постоянной задержки с помощью подходящего фазового уравнивателя, однако оно не может быть устранено. Так, передавая сигнал через неминимально-фазовый фильтр, можно восстановить его только после задержки, т. е. получится $f(t-\alpha)$ с положительным α .

4. Основное выражение для средне-квадратичной ошибки

Предположим, что в качестве предсказывающе-сглаживающего фильтра на рис. 1 применен фильтр с частотной характеристикой $Y(\omega)$. Какова средне-квадратичная ошибка предсказания? Так как различные частоты некогерентны, можно вычислить среднюю мощность ошибки

$$e(t) = s(t + \alpha) - g(t) \quad (9)$$

суммированием составляющих различных частот. Рассмотрим компоненты сигнала и шума на определенной частоте ω_1 . Предполагается, что сигнал и шум некогерентны при всех частотах. Тогда на частоте ω_1 составляющая ошибки, обусловленная шумом, равна $N(\omega_1) |Y(\omega_1)|^2$, где $N(\omega_1)$ — средняя мощность шума на частоте ω_1 .

Существует также составляющая ошибки благодаря ослаблению компонент сигнала после прохождения его через фильтр. Для составляющей частоты ω_1 амплитуды на входе и на выходе должны быть равны, а фаза на выходе должна иметь опережение на $\alpha\omega_1$. Тогда ошибка в мощности будет

$$|Y(\omega_1) - e^{i\alpha\omega_1}|^2 P(\omega_1), \quad (10)$$

где $P(\omega_1)$ — мощность сигнала на частоте ω_1 .

Общая средне-квадратичная ошибка на частоте ω_1 является суммой этих двух ошибок, или

$$E_{\omega_1} = |Y(\omega_1)|^2 N(\omega_1) + |Y(\omega_1) - e^{i\omega_1\alpha}|^2 P(\omega_1), \quad (11)$$

и общая средне-квадратичная ошибка для всех частот равна

$$E = \int_{-\infty}^{\infty} [|Y(\omega)|^2 N(\omega) + |Y(\omega) - e^{i\alpha\omega}|^2 P(\omega)] d\omega. \quad (12)$$

Задача заключается в минимизации E надлежащим выбором $Y(\omega)$ с учетом того, что система с частотной характеристикой $Y(\omega)$ должна быть физически осуществима.

Некоторые важные выводы можно сделать сразу же из рассмотрения выражения (12). Сигнал и шум входят в уравнение только

через посредство их спектров мощности. Следовательно, единственные статистические параметры сигнала и шума, необходимые для решения проблемы, это их спектры. Два различных сигнала с одним и тем же спектром мощности ведут к одному и тому же оптимальному предсказывающему фильтру и к той же самой средне-квадратичной ошибке. Например, если в качестве сигнала используется речь, она может быть предсказана тем же фильтром, который используется для предсказания белого теплового шума, предварительно пропущенного через фильтр, дающий на выходе такой же спектр, как спектр речи.

Говоря более вольно, линейный фильтр может использовать статистику, относящуюся только к амплитудам различных частотных составляющих, статистика фазовых углов этих составляющих не может быть использована. Только нелинейное предсказание может использовать этот статистический эффект для улучшения предсказания.

Ясно, что в проблеме линейной аппроксимации по методу наименьших квадратов можно при желании заменить сигнал и шум любыми временными последовательностями, которые имеют те же спектры мощности. Это никоим образом не изменит оптимального фильтра и средне-квадратичной ошибки.

5. Чистая проблема сглаживания

Основная трудность в минимизации выражения (12) для средне-квадратичной ошибки лежит в надлежащем введении условия, что $Y(\omega)$ должна быть физически осуществимой частотной характеристикой. Сначала решим вопрос без этого ограничения, а затем из этого решения построим лучший физически осуществимый фильтр.

Отказ от условия физической осуществимости равносильно допущению произвольного выбора $Y(\omega)$ или, что то же, любой импульсной реакции $K(t)$. Так, $K(t)$ не обязательно равна нулю для $t < 0$ и допускается весовая функция, которая может быть применена как к прошлому, так и к будущему $f(t)$. Другими словами, предположим, что при предсказании может быть использована полная функция $f(t) = s(t) + n(t)$ при t , меняющемся от $-\infty$ до $+\infty$.

Предположим, что в выражении (12)

$$Y(\omega) = C(\omega) e^{iB(\omega)} \quad (13)$$

с действительными $C(\omega)$ и $B(\omega)$.

Тогда получим

$$E = \int_{-\infty}^{\infty} [C^2 N + P(C^2 + 1 - 2C \cos(\alpha\omega - B))] d\omega, \quad (14)$$

где $C(\omega)$, $N(\omega)$ и т. д. будем и далее для простоты писать как C , N и т. д. Ясно, что наилучшим выбором $B(\omega)$ будет $B(\omega) = a\omega$, так как это максимизирует $\cos[a\omega - B(\omega)]$. Тогда получим из (14)

$$E = \int_{-\infty}^{\infty} [C^2(P+N) - 2PC + P] d\omega. \quad (15)$$

Дополняя равенство (15) до квадрата добавлением и вычитанием $P^2/(P+N)$, получим

$$E = \int_{-\infty}^{\infty} \left[C^2(P+N) - 2PC + \frac{P^2}{P+N} - \frac{P^2}{P+N} + P \right] d\omega \quad (16)$$

или

$$E = \int_{-\infty}^{\infty} \left(\left[\sqrt{P+N}C - \frac{P}{\sqrt{P+N}} \right]^2 + \frac{PN}{P+N} \right) d\omega. \quad (17)$$

Член в квадратных скобках является квадратом действительного числа, поэтому он положителен или равен нулю. Ясно, что для минимизации E надо выбрать C таким, чтобы этот член всегда был равен нулю, т. е.

$$C(\omega) = \frac{P(\omega)}{P(\omega) + N(\omega)},$$

$$Y(\omega) = \frac{P(\omega)}{P(\omega) + N(\omega)} e^{i\alpha\omega}. \quad (18)$$

При таком выборе $Y(\omega)$ средне-квадратичная ошибка (17) будет равна

$$E = \int_{-\infty}^{\infty} \frac{P(\omega)N(\omega)}{P(\omega) + N(\omega)} d\omega. \quad (19)$$

Наилучшая весовая функция дается обратным преобразованием Фурье выражения (18)

$$K(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \frac{P(\omega)}{P(\omega) + N(\omega)} e^{i\omega(t+\alpha)} d\omega. \quad (20)$$

$K(t)$ будет, вообще говоря, меняться от $t = -\infty$ до $t = \infty$. Она не является импульсной реакцией физически осуществимого фильтра. Однако это вполне подходящая весовая функция. Если бы можно было ждать, пока вся функция $s(t) + n(t)$ будет доступна, то эта весовая функция была бы как раз той, которая требуется для определения $s(t + \alpha)$.

Иначе говоря, весовая функция $K(\tau)$ может быть получена для физически осуществимого фильтра, если допускается достаточная задержка, так что $K(\tau)$ имеет практически нулевое значение для всего будущего. Таким образом, решена проблема сглаживания с «бесконечным сдвигом». Хотя $Y(\omega)$ в формулах (18) физически неосуществима, $Y(\omega)e^{-i\beta\omega}$ будет осуществима или почти осуществима, если β взято достаточно большим.

6. Проблема чистого предсказания

Теперь рассмотрим другой специальный случай — отсутствие искажающего шума. Это чистое предсказание того, какова лучшая оценка $s(t+\alpha)$, когда известно $s(t)$ от $t=-\infty$ до 0?

Было показано, что решение должно зависеть только от спектра мощности сигнала и шума, и так как предполагается, что шум тождественно равен нулю, решение зависит только от спектра мощности $P(\omega)$ сигнала. Если это так, то можно заменить действительный сигнал любым другим, имеющим тот же самый спектр; решение задачи о наилучшем предсказывающем фильтре останется для измененной таким образом задачи прежним.

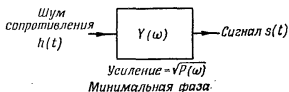
Любой желаемый спектр $P(\omega)$ может быть получен при прохождении широкополосного шума сопротивления или белого шума через формирующий фильтр, характеристика усиления которого есть $\sqrt{P(\omega)}$. Спектр шума сопротивления плоский (по крайней мере вплоть до частот, высших, чем любые частоты, имеющие применение в связи), и фильтр просто умножает этот постоянный спектр на квадрат усиления фильтра $P(\omega)$. Фазовая характеристика этого фильтра может быть выбрана любым способом, согласующимся с условием физической осуществимости. Выберем фазовую характеристику так, чтобы фильтр был минимально-фазовым для усиления $\sqrt{P(\omega)}$. Тогда фильтр будет иметь фазовую характеристику, определяемую как

$$B(\omega_0) = \frac{-\omega_0}{\pi} \int_0^{\infty} \frac{\log P(\omega) - \log P(\omega_0)}{\omega^2 - \omega_0^2} d\omega. \quad (21)$$

Более того, этот минимально-фазовый фильтр имеет физически осуществимый обратный.

Задача сведена к форме, показанной на рис. 5. В нашем распоряжении имеется функция $s(t)$ вплоть до $t=0$. Это эквивалентно знанию шума сопротивления $h(t)$ до $t=0$, так как фильтр Y имеет физически осуществимое обращение и можно пропустить имеющийся сигнал $s(t)$ через обращение Y^{-1} для получения $h(t)$.

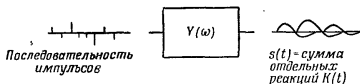
Поэтому задача формулируется следующим образом: как наилучшим способом обработать $h(t)$, чтобы аппроксимировать $s(t+\alpha)$ в смысле наименьших квадратов. Эта задача легко решается. Шум сопротивления можно представлять себе как совокупность большого числа близко расположенных и очень коротких импульсов, как показано на рис. 6. Импульсы имеют гауссовское распределение амплитуд и статистически независимы. Каждый из этих импульсов,



Р и с. 5. Построение спектра сигнала из белого шума.

поступая на фильтр Y , образует реакцию, соответствующий импульсной реакции фильтра, как показано на рис. 6, и сигнал $s(t)$ является суммой этих элементарных откликов.

Что нам известно, так это $h(t)$ вплоть до данного момента, т. е. эффективно известны импульсы до $t=0$ и ничего не известно о них после $t=0$ (они еще не возникли). Будущий сигнал $s(t+\alpha)$, таким



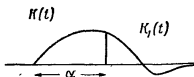
Р и с. 6. Результат «воздействия» белого шума на входе.

образом, состоит из двух частей: хвостов реакций на импульсы, которые уже прошли, и части, обусловленной теми импульсами, которые должны поступить на фильтр за время от $t=0$ до $t=\alpha$. Первая часть полностью предсказуема, тогда как вторая часть полностью не предсказуема, будучи статистически независима от имеющейся у нас к текущему моменту информации.

Общий результат первой части может быть получен построением фильтра, импульсная реакция которого является хвостом импульсной реакции фильтра Y , сдвинутой вперед на α секунд. Это показано на рис. 7, где $K_1(t)$ является новой импульсной реакцией, а $K(t)$ — старой. Новый фильтр реагирует на импульс, появившийся в данный момент, так, как фильтр Y будет реагировать через α секунд. Если $h(t)$ используется как входной сигнал для

этого нового фильтра Y_1 , то реакция будет теперь предсказуемой частью будущей реакции Y на тот же входной сигнал через α секунд.

Вторая, или непредсказуемая часть будущей реакции, соответствующая импульсам, которые еще не появились, конечно, не может быть построена. Однако известно, что средняя величина этой части должна быть нулем, так как будущие импульсы с одинаковой вероятностью могут иметь как один знак, так и противоположный.



Р и с. 7. Построение физически осуществимой реакции $K_1(t)$ по $K(t)$.

Итак, арифметическое среднее, или центр тяжести возможных будущих реакций, есть предсказуемая часть, даваемая реакцией Y_1 . Но, как известно, арифметическое среднее любого распределения есть точка, относительно которой средне-квадратичное отклонение наименьшее. Таким образом, реакция Y_1 является желаемым предсказанием $s(t+\alpha)$.

При построении Y_1 предполагалось, что имеется в распоряжении белый шум $h(t)$. В действительности же данным является сигнал $s(t)$. Следовательно, наилучшей операцией над имеющимися данными является операция $Y_1(\omega) Y^{-1}(\omega)$; множитель $Y^{-1}(\omega)$ преобразует функцию $s(t)$ в белый шум $h(t)$, а второй множитель $Y_1(\omega)$ производит наилучшее предсказание, основанное на $h(t)$.

Решение вопроса может быть суммировано следующим образом.

1. Определим минимально-фазовый фильтр, имеющий характеристику усиления $\sqrt{P}(\omega)$. Пусть комплексная частотная характеристика этого фильтра будет $Y(\omega)$, а его импульсная реакция $K(t)$.

2. Построим фильтр, импульсная реакция которого

$$K_1(t) = \begin{cases} K(t+\alpha) & \text{для } t \geq 0, \\ 0 & \text{для } t < 0. \end{cases} \quad (22)$$

Пусть частотная характеристика этого фильтра есть $Y_1(\omega)$.

3. Оптимальный в смысле наименьших квадратов предсказывающий фильтр имеет тогда характеристику

$$Y_1(\omega) Y^{-1}(\omega). \quad (23)$$

Средне-квадратичная ошибка E предсказания легко вычисляется. Ошибка возникает из-за импульсов, появляющихся во время от

$t=0$ до $t=\alpha$. Так как эти импульсы не коррелированы, средне-квадратичная сумма ошибок равна сумме отдельных средне-квадратичных ошибок. Отдельные импульсы действуют при образовании средне-квадратичной ошибки пропорционально квадрату $K(\alpha-t)$. Отсюда общая средне-квадратичная ошибка дается выражением

$$E^2 = \varrho \int_0^{\alpha} K^2(\alpha-t) dt = \varrho \int_0^{\alpha} K^2(t) dt, \quad (24)$$

где $\varrho = \int p(\omega) d\omega$ — средний квадрат сигнала. На том же основании средний квадрат функции $s(t+\alpha)$ равен

$$U^2 = \varrho \int_0^{\infty} K^2(t) dt, \quad (25)$$

и относительная ошибка предсказания равна отношению средне-квадратичной ошибки к средне-квадратичному значению $s(t+\alpha)$, т. е.

$$\frac{E}{U} = \left[\frac{\int_0^{\alpha} K^2(t) dt}{\int_0^{\infty} K^2(t) dt} \right]^{\frac{1}{2}}. \quad (26)$$

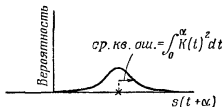
Предсказание будет относительно плохим, если площадь под кривой $K^2(t)$ вплоть до α велика по сравнению со всей площадью, и хорошим, если она мала по сравнению со всей площадью. Очевидно из (26), что относительная ошибка начинается с нуля для $\alpha=0$ и монотонно возрастает с α , стремясь к единице при $\alpha \rightarrow \infty$.

Существует один важный случай, в котором приведенные сообщения могут дать много больше. В нашем анализе истинная задача была заменена задачей, в которой сигнал является временным рядом гауссовского типа, полученным путем пропускания шума сопротивления через фильтр с усилением $\sqrt{P(\omega)}$. Предположим, что исходный сигнал уже является временной последовательностью этого типа. Тогда ошибка в предсказании, связанная с хвостами импульсов, появляющихся между $t=0$ и $t=\alpha$, будет иметь гауссовское распределение. Это следует из того, что каждый импульс имеет гауссовское распределение, а сумма любого числа гауссовских величин будет также гауссовской величиной. Стандартное отклонение этого распределения ошибок является как раз средне-квадратичной ошибкой E , полученной из равенства (24).

Другими словами, на основе имеющихся данных об $s(t)$ для $t < 0$ можно сказать, что будущее значение сигнала $s(t+\alpha)$ имеет

гауссовское распределение. Для предсказуемой величины наилучший линейный предсказатель выделит центр этого распределения. Истинное значение будущей величины будет отличаться от него, как показано на рис. 8, где будущее значение функции нанесено по горизонтали, а плотность вероятности для различных величин $s(t+\alpha)$ — по вертикали.

Ясно, что в этом частном случае метод линейного предсказания является в известном смысле наилучшим из возможных. Центр гауссовского распределения останется естественной выборочной точкой, даже если заменить средне-квадратичный критерий любым



Р и с. 8. Распределение ошибок предсказания в случае гауссовского распределения.

другим разумным критерием, скажем медианой, наивероятнейшим значением и т. п. Нелинейное предсказание в этом случае ничего лучшего не дает по сравнению с линейным. В общем случае, с другой стороны, распределение будущих значений не будет гауссовским и форма кривой распределения может меняться от точки к точке в зависимости от частной предыстории кривой. В этом случае нелинейная схема может дать лучшие результаты по сравнению с линейной; точные характеристики оптимальной операции будут существенно зависеть от принятого критерия для наилучшего предсказания.

7. Предсказание при наличии шума

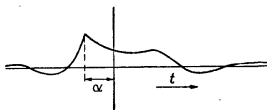
Рассмотрим теперь общую проблему предсказания и сглаживания при наличии шума. Требуется найти наилучшую оценку для $s(t+\alpha)$, когда известна функция $s(t)+n(t)$ от $t=-\infty$ до настоящего момента времени. Если сигнал $s(t)+n(t)$ пропущен через фильтр, усиление которого $[P(\omega)+N(\omega)]^{-1/2}$, то получится плоский спектр, который можно трактовать как белый шум. Пусть $Y_1(\omega)$ — характеристика минимально-фазового фильтра, имеющего такое усиление. Тогда как $Y_1(\omega)$, так и обращение $Y_1^{-1}(\omega)$ являются физически осуществимыми фильтрами. Очевидно, значение сигнала на входе Y_1 и знание его реакции эквивалентны. Наилучшая линейная операция над реакцией будет давать то же самое предсказание, как и соответствующая наилучшая линейная операция на входе.

Если бы была известна полная функция $s(t) + n(t)$ от $t = -\infty$ до $t = \infty$, то наилучшей операцией, примененной к входу, была бы та, которая удовлетворяет равенству (18). Если в качестве фазовой характеристики взято $B(\omega)$, то это эквивалентно операции

$$Y_2(\omega) = \frac{P(\omega)}{[P(\omega) + N(\omega)]^{1/2}} e^{i[\alpha\omega - B(\omega)]} \quad (27)$$

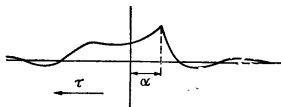
над реакцией Y_1 , имеющей характер белого шума.

Пусть импульсная реакция, полученная из (27), будет $K_2(t)$, как показано на рис. 9. $K_2(t)$ в общем случае содержит хвосты, простирающиеся как до $t = -\infty$, так и до $t = \infty$, стык между двумя половинами кривой смещен из начала координат на время предсказания α .



Р и с. 9. Возможная функция $K_2(t)$.

Функция $K_2(\tau)$ на рис. 10, конечно, идеальная весовая функция в применении к белому шуму на выходе Y_1 . Однако единственные фактически имеющиеся в нашем распоряжении данные



Р и с. 10. Весовая функция $K_2(\tau)$, соответствующая рис. 9.

к моменту $\tau=0$ — это импульсы, которые, можно полагать, появились в прошедшей истории этого выходного процесса. Какой вес следует придать этим импульсам, чтобы получить лучшее предсказание? Кажется естественным взвесить их так, как если бы все данные были уже налицо, и приписать нулевой вес будущим импульсам (ведь требуется сделать фильтр физически осуществимым). Что это действительно правильное взвешивание, когда различные входные импульсы статистически независимы, будет сейчас выведено как следствие общих статистических принципов.

Предположим, что имеются случайные переменные $x_1, x_2, \dots, x_n$, которые статистически независимы или по крайней мере обладают тем свойством, что среднее произведение любых двух значений $\overline{x_m x_n}$ равно нулю. Эти переменные будем трактовать как амплитуды отдельных импульсов шума и к ним попробуем применить весовую функцию, изображенную на рис. 10.

Пусть y есть другая случайная величина, коррелированная с $x_1, x_2, \dots, x_n$, оцениваемая в смысле наименьших квадратов при помощи линейной операции над $x_1, x_2, \dots, x_n$. В нашем случае под y будем понимать сигнал $s(t)$ спустя α секунд.

Предсказываемая величина будет иметь вид

$$y_1 = \sum_{i=1}^n a_i x_i$$

и средне-квадратичная ошибка будет равна

$$\begin{aligned} E &= \overline{(y - y_1)^2} = \overline{(y - \sum (a_i x_i))^2} = \\ &= \overline{y^2} - 2 \sum_{i=1}^n a_i \overline{x_i y} + \sum_{i,j=1}^n a_i a_j \overline{x_i x_j} = \\ &= \overline{y^2} - 2 \sum a_i \overline{x_i y} + \sum_{i=1}^n a_i^2 \overline{x_i^2}, \end{aligned} \quad (28)$$

так как все слагаемые в двойной сумме пропадают, за исключением тех, для которых $i = j$. Найдем минимальное E подходящим выбором a_i . Приравняв частные производные по a_i нулю, получим

$$\frac{\partial E}{\partial a_i} = -2 \overline{x_i y} + 2 a_i \overline{x_i^2} = 0,$$

или

$$a_i = \frac{\overline{x_i y}}{\overline{x_i^2}}. \quad (29)$$

В вышеприведенном вычислении существенно то, что каждое из n минимизирующих уравнений содержит только одно a_i ; $\frac{\partial E}{\partial a_i}$ содержит только a_i и т. д. Другими словами, минимизация E по всем a_i эквивалентна минимизации раздельно по каждому a_i ; a_i будет иметь величину $\overline{x_i y} / \overline{x_i^2}$, каковы бы ни были значения, приписанные другим a_i .

Возвратимся теперь к проблеме предсказания и сглаживания. Функция $K_2(\tau)$ дает надлежащий вес импульсам, если все их можно использовать. Условие физической осуществимости требует, чтобы

будущие импульсы, соответствующие $\tau < 0$, давались с нулевым весом. На основании вышеприведенного статистического принципа импульсам, поступившим в прошлом, по-прежнему придается вес $K_2(\tau)$. Другими словами, надлежащий фильтр для белого шума на входе имеет импульсную реакцию, равную нулю для $t < 0$ и $K_2(t)$ для $t > 0$.

Суммируя, получим следующие этапы решения.

1. Вычисляется минимально-фазовая частотная характеристика для усиления $(P+N)^{-1/2}$; обозначим ее $Y_1(\omega)$.

2. Пусть имеется равенство:

$$Y_2(\omega) = Y_1^{-1}(\omega) \frac{P(\omega)}{N(\omega) + P(\omega)}.$$

Это физически неосуществимая характеристика. Пусть ее преобразование Фурье есть $K_2(t)$.

3. Положим $K_3(t) = K_2(t + \alpha)$ для $t \geq 0$ и $K_3(t) = 0$ для $t < 0$. Переход от K_2 к K_3 означает отсечение первых α секунд и сдвиг остающегося хвоста до $t = 0$. Эта импульсная реакция физически реализуемого фильтра и является оптимальной операцией над прошлым входного белого шума. Пусть соответствующая функция передачи обозначена $Y_3(\omega)$.

4. Построим $Y_4(\omega) = Y_3(\omega) Y_1(\omega)$. Это — оптимальный сглаживающий и предсказывающий фильтр для сигнала $s(t) + n(t)$.

Как и в проблеме чистого предсказания, если шум и сигнал оказываются чисто гауссовскими временными последовательностями, то линейное предсказание является абсолютно оптимальным среди всех операций предсказания, линейных или нелинейных. Более того, распределение значений $s(t + \alpha)$, когда $f(t)$ известно для $t < 0$, является гауссовским.

8. Обобщения

Наша теория может быть обобщена в нескольких направлениях. Эти обобщения будут упомянуты только кратко, однако могут быть получены методами, сходными с рассмотренными выше.

Во-первых, предполагается, что истинный сигнал и шум не коррелированы. Относительно простое обобщение понятий, используемых в разд. 4, позволяет учесть корреляцию между этими временными последовательностями.

Второе обобщение относится к случаю, когда имеется несколько коррелированных временных последовательностей, скажем $f_1(t), f_2(t), \dots, f_n(t)$. Желательно предсказать $s_1(t + \alpha)$ на основе знания $f_1, f_2, \dots, f_n$.

Наконец, искомой величиной может быть не $s(t + \alpha)$, а например $s'(t + \alpha)$ — будущее значение производной истинного сигнала. В этом

случае проблема может быть эффективно сведена к уже решенной проблеме простым взятием производной. Функция $f(t)$ пропускается через дифференцирующее устройство для образования $g(t) = f'(t)$. Тогда определено наилучшее линейное предсказание для $g(t)$.

9. Обсуждение основных предположений

Результат в прикладной математике надежен постольку, поскольку надежны предположения, из которых он выведен. Развита выше теория особенно часто приводит к попыткам ее неудачных применений, так как трудно решить, являются ли в данном частном случае основные предположения удовлетворительным описанием физической ситуации. Тот, кто использует эту теорию, должен тщательно проанализировать каждое из трех предположений как в теории сглаживания, так и в теории предсказания.

Предположение, что сигнал и шум являются стационарными, вероятно, наиболее безобидное из всех трех, так как из общего характера задачи очевидно место, когда это предположение нарушается. При определении требуемого спектра мощности $P(\omega)$ и $N(\omega)$ часто обнаруживаются какие-либо временные вариации статистической структуры временных рядов. Если эти вариации медленны по сравнению с другими временными постоянными, такие нестационарные вопросы все же могут решаться квазистационарными методами. Может быть спроектирован линейный предсказатель, характеристика которого медленно изменяется так, что она остается оптимальной для «локальной» статистики.

Понятие средне-квадратичного критерия более сложно, так как здесь включаются не только качественные, но и количественные вопросы. При минимизации средне-квадратичной ошибки в действительности внимание уделяется очень большим ошибкам. В целом предсказание выбирается так, чтобы сделать эти ошибки столь малыми, сколь возможно без особой заботы об относительно малых ошибках. Однако во многих случаях важно делать по возможности чаще очень точные предсказания, даже если за счет этого будут возникать отдельные большие ошибки. Когда распределение будущих событий гауссовское, безразлично, какой критерий применяется, так как наиболее вероятное событие совпадает с тем, для которого средне-квадратичная ошибка наименьшая. Однако при асимметричном или многовершинном распределении вопрос приобретает актуальность.

Как пример, рассмотрим предсказание погоды: будет ли завтра ясный день? Так как ясных дней большинство и не существует дней с отрицательными осадками, дополняющими дождливые дни, здесь распределение асимметрично. Для такого распределения средняя точка, получаемая в результате предсказания, мини-

мизирующего средне-квадратичную ошибку, может быть представлена как день с небольшим дождиком. Однако для человека, собирающегося на пикник, такое предсказание не будет иметь цены. Его интересует вероятность действительно ясного дня. Если пикник должен быть отменен из-за любого дождя, то количество осадков имеет относительно малое значение.

В качестве второго примера рассмотрим задачу перехвата автомобиля с преступниками, пытающимися спастись бегством по сети дорог. Если сразу впереди на дороге встречается развилка, ясно, что преследователь должен расположиться либо на одной ветви, либо на другой, делая этот выбор, если он необходим, случайным образом. Средне-квадратичная ошибка перехвата будет, однако, минимальной, если он расположится в поле за развилкой. Такие же проблемы могут возникнуть в артиллерии, когда обычно интересуются числом действительных попаданий независимо от числа промахов.

Третье предположение — условие линейности нельзя назвать ни качественным, ни количественным. Это скорее добровольное ограничение типа операций или устройств, применяемых для предсказания. Математические основания этого предположения ясны: линейные проблемы всегда более просты, чем их нелинейные обобщения. В некоторых применениях предположение линейности может быть обосновано одним из следующих соображений.

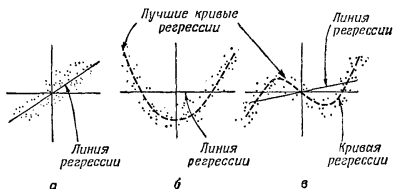
1. Линейный предсказатель может быть абсолютно оптимальным в упомянутом случае предсказания гауссовских временных рядов.

2. Линейное предсказание может диктоваться требованиями простоты реализации. Линейные фильтры легко синтезировать, и существует развитая теория их построения, чего нельзя сказать о нелинейных фильтрах.

3. Можно применить линейную теорию просто из-за отсутствия другой теории. Неполное решение лучше, чем отсутствие решения вообще.

Что теряется при ограничении линейным предсказанием? То, что нелинейные эффекты могут быть значительными при предсказании, можно проиллюстрировать, вернувшись к примеру предсказания завтрашней погоды. Известно, что для определения будущего, важнее цепь событий на протяжении некоторого времени, чем отдельно взятые случаи. Например, при прохождении холодного или теплого фронта характерной является определенная последовательность явлений. Более того, значение данного события может зависеть в значительной степени от интенсивности, с которой оно происходит. Так, резкое падение барометра означает, что наступает довольно неприятная погода. Вдвое большее падение за то же время, с другой стороны, не означает просто, что погода станет вдвое более неприятной, оно может предсказывать ураган.

В заключение отметим, что требования, чтобы предсказание было линейным и в то же время минимизировало средне-квадратичную ошибку, не всегда совместимы. Абсолютно наилучшее средне-квадратичное предсказание (игнорируя предположение о линейности), конечно, всегда укажет среднее будущего распределения, т. е. «центр тяжести», так как в любом случае это минимизирует средне-квадратичную ошибку. В общем случае, однако, положение этого центра тяжести будет нелинейной функцией прошлых событий. Когда требуется, чтобы предсказание было *линейной* операцией над прошлыми событиями, математик вынужден идти на компромисс между противоречивыми требованиями этих прошлых событий. Компромисс равнозначен по существу усреднению по всем



Р и с. 11. Некоторые диаграммы рассеяния с линиями и кривыми регрессии.

фазовым значениям различных составляющих сигнала; любая уместная информация, имеющаяся в фазовых значениях, не может быть использована.

Это может быть показано на примере известной статистической задачи о вычислении линии или плоскости регрессии для обеспечения линейной оценки переменной y по ряду известных переменных, коррелированных с y методом наименьших квадратов¹⁾. Самый простой случай тот, когда есть только одна переменная x и одна неизвестная y , оцениваемая по x . На рис. 11 приведены три «диаграммы рассеяния», используемые в статистике. Переменной x может быть, например, вес человека, а y — его рост. Большое количество значений нанесено на рис. 11. Желательно оценить или предсказать рост человека, зная только его вес. Если условиться использовать только линейные операции, y должен вычисляться по формуле $y=ax$. Наилучшим выбором a для минимально-квд-

¹⁾ Noel P., Introduction to mathematical statistics, Wiley, N. Y., 1947.

ратичного предсказания является $\overline{xy}/\overline{x^2}$, и соответствующая прямая линия известна как линия регрессии. Случай нормального распределения соответствует шуму гауссовского типа, в котором линейное предсказание является абсолютно оптимальным.

Рис. 11б и 11, в представляют собой диаграммы рассеяния для других распределений двух переменных. Линии регрессии теперь не являются таким хорошим предсказанием для y , как на рис. 11, а. Условие, чтобы предсказываемая величина была *линейной* функцией известных данных, требует компромисса, который может быть очень серьезным. Очевидно из рис. 11б и 11, в, что значительно лучшая оценка для y может быть получена, если допустить нелинейные операции над x . В частности, функции ax^2+b и cx^3+dx — более подходящие.

Для предсказания y по двум известным переменным x_1 и x_2 можно построить диаграмму рассеяния в трех измерениях. Линейное предсказание требует проведения плоскости регрессии через систему точек. Если имеется n известных величин $x_1, x_2, \dots, x_n$, то необходимо пространство $n+1$ измерений и линейная теория соответствует нахождению гиперплоскости n измерений.

Проблема сглаживания и предсказания для временных последовательностей аналогична. Однако при этом мы имеем дело с функциональным пространством, определенным всеми значениями $f(t)$ при $t < 0$. Оптимальное линейное предсказание соответствует гиперплоскости в этом функциональном пространстве.

МАТЕМАТИЧЕСКАЯ ТЕОРИЯ ДИФФЕРЕНЦИАЛЬНОГО АНАЛИЗАТОРА¹⁾

Введение

Дифференциальный анализатор — это машина, разработанная в Массачусетском технологическом институте под руководством доктора В. Буша для получения численных решений обыкновенных дифференциальных уравнений²⁾. Основные принципы, лежащие в основе дифференциального анализатора, впервые были сформулированы Кельвином, однако в то время вследствие технических трудностей было невозможно построить машину задуманного им типа. Доктор Буш и его сотрудники независимо открыли те же самые принципы, и в 1931 г. было завершено создание первого дифференциального анализатора. Технические трудности, возникшие при создании машины, были преодолены с помощью нескольких остроумных устройств, как, например, усилители напряжений вращения и блоки корректировки мертвого хода, а также посредством улучшенных технических приемов работы машины. С тех пор в различных частях света было построено несколько других машин. Они использовались при решении многих задач, возникающих в технике, физике и других отраслях науки.

Механическая работа машины и способы ее подготовки к работе подробно описаны в следующих статьях³⁾. Для наших целей можно кратко резюмировать способ работы машины следующим образом. Каждый член или переменное данного уравнения или данной системы уравнений представляется в машине определенным валом. В любой момент значение переменной пропорционально

¹⁾ Shannon C., *Mathematical theory of the differential analyzer*, *J. Math. and Phys.*, 20, № 4 (1941), 337.

²⁾ Аналогичный прибор раньше был разработан А. Н. Крыловым. Этот прибор испытывался и улучшался в 1911—1914 гг. Описание опубликовано в *Изв. АН*, 5 сер., 20, № 1 (1904) 17—37. (См. также Крылов А. Н., *Собрание трудов*, т. V, 1937, 547—574). —Прим. ред.

³⁾ Bush V., *The differential analyzer, a new mashine for solving differential equations*, *Journ. Franklin Institute*, 212 (1931), 447. Hartree D., *The mechanical integration of differential equations*, *Mathematical gazette*, XXII (1938), 342.

числу полных оборотов, совершенных соответствующим валом из некоторого фиксированного положения. Эти валы связаны между собой с помощью механических устройств, обеспечивающих выполнение определенных математических соотношений между числами оборотов взаимно связанных валов. Наиболее важными механическими устройствами являются устройства четырех типов¹⁾: коробки скоростей, сумматоры, интеграторы и функциональные устройства. С помощью этих устройств для соответствующих валов обеспечивается выполнение любого соотношения, определяемого данными уравнениями. Например, если в дифференциальном уравнении сумма двух членов равна третьему члену, то соответствующие валы анализатора связываются суммирующим устройством. Если в уравнении появляются члены x , y и dy/dx , то для осуществления соотношения $y = \int (dy/dx \, dx)$ соответствующие валы соединяются интегратором и т. д.

Вращение вала, представляющего независимое переменное, влечет за собой вращение всех других валов в соответствии с уравнением. Таким образом, численное решение может быть получено подсчетом чисел оборотов вала, представляющего зависимое переменное, соответствующих равным приращениям независимого переменного, причем результат вычерчивается в виде кривой. С помощью специального устройства вывода можно заставить машину автоматически вычерчивать кривые решения уравнений.

Когда дифференциальный анализатор монтировался впервые, предполагалось, что все функциональные соотношения между членами решаемого уравнения будут вводиться в машину с помощью функциональных устройств. Однако в связи с одной баллистической задачей, в которой требуется вычисление функции x^2 , было замечено, что эта функция может быть получена без функционального устройства, а с помощью присоединения интегратора для выполнения операции $2 \int_0^x x \, dx$. Вскоре было обнаружено, что практически все важные простые функции могут быть реализованы путем приме-

¹⁾ Разработаны и некоторые другие специальные устройства. Дифференциальный анализатор Массачусетского технологического института имеет множительное устройство, предназначенное для получения произведения двух членов, а дифференциальный анализатор в Манчестере имеет устройство, пригодное для решения смешанных разностно-дифференциальных уравнений. Однако эти устройства значительно менее существенны, чем устройства, упомянутые выше; как будет показано позднее, множительное устройство в действительности всегда может быть заменено двумя интеграторами и сумматором.

нения одних только интеграторов и сумматоров. Это осуществляется посредством установки на анализаторе вспомогательного уравнения, решением которого является требуемая функция. Дж. Герьери¹⁾ в диссертации, написанной в 1932 г., устанавливает соотношения для реализации большинства элементарных функций.

В данной статье рассматриваются математические вопросы теории дифференциального анализатора. Наиболее важные результаты связаны с условиями, при которых могут быть реализованы функции от одного или нескольких переменных, а также условиями, при которых возможно решение обыкновенных дифференциальных уравнений. Кроме того, уделяется внимание аппроксимации функций (которые не могут быть реализованы точно), аппроксимации передаточных чисел и автоматическому управлению скоростью работы устройства.

В статье предполагается, что все рассматриваемые обыкновенные дифференциальные уравнения имеют единственные решения и что все формальные процессы дифференцирования, интегрирования и т. д. правомерны в рассматриваемой области изменения аргументов. В случае общих дифференциальных уравнений нет необходимости, чтобы уравнения были интегрируемы; однако предполагается, что решение существует вдоль любой кривой в данной области. Основания для этого будут введены ниже.

Будем рассматривать идеализированный дифференциальный анализатор в предположении, что в нашем распоряжении имеется неограниченное число следующих идеальных блоков.

1. *Интеграторы.* Для данных двух валов u и v интегратор вынуждает третий вал w вращаться при всех изменениях u и v в соответствии с соотношением $w = \int_{v_0}^v (u+a)dv$, где a — произвольная

константа. В реальных интеграторах максимальное значение $|u+a|$ ограничено, но может быть сделано сколь угодно большим посредством изменения масштабных множителей, с тем чтобы интегрирование могло быть выполнено везде, за исключением полюсов функции u . Константа a представляет собой начальную установку интегратора.

2. *Сумматоры.* Для данных двух валов u и v сумматор вынуждает третий вал w вращаться, реализуя $u+v$ при всех изменениях u и v . Если не учитывать мертвого хода и люфта шестерен, то для реальных дифференциальных суммирующих устройств эти условия выполняются. Легко видеть, что любой способ соединения сум-

¹⁾ Guerrieri J., Methods of introducing functional relations automatically on the differential analyser, S. M. Thesis, M. I. T., 1932.

маторов со свободными валами $X_1, \dots, X_n$, Y осуществляет соотношение $Y = \sum_{k=1}^n a_k X_k$, где a_k — постоянные действительные числа.

Для удобства будем называть такое соединение сумматором. Соединив последовательно простые сумматоры, можно сделать все значения a_k единичными и получить $Y = \sum_{k=1}^n X_k$.

3. *Функциональные устройства.* Для данного вала x функциональное устройство обеспечивает вращение вала y в соответствии с заданным соотношением $y = f(x)$, где $f(x)$ — некоторая функция, имеющая лишь конечное число конечных разрывов.

4. *Коробки скоростей.* Если дан вал x , то коробка скоростей с передаточным числом k создает для второго вала режим вращения kx . Тогда, сделав величину u тождественно равной нулю, получим на интеграторе $\omega = av$. Очевидно, что коробки скоростей являются теоретически излишними и применяются для экономии. Поэтому рассматривать коробки скоростей не будем, а только покажем, что любое передаточное число k может быть аппроксимировано с применением пар шестерен только двух размеров.

Наконец, предполагается, что каждый из описанных выше элементов способен выполнять свою частную задачу тогда и только тогда, когда на каждый вал действует не более чем один источник движения. Под источником движения понимается любой из следующих: вал независимого переменного, вал ω интегратора, вал ω сумматора (или Y в общем случае), вал $f(x)$ функционального устройства или вал kx коробки скоростей. Это условие чрезвычайно важно при установке устройств дифференциального анализатора. Оно накладывает ограничение на возможные соединения блоков и лежит в основе настоящего анализа.

Будем говорить, что можно решить систему обыкновенных дифференциальных уравнений с независимым переменным x и зависимыми переменными $y_1, \dots, y_n$ тогда и только тогда, когда может быть найдено устройство, использующее перечисленные выше элементы и удовлетворяющее предположению об источнике движения, и такое, что вращение вала независимого переменного x влечет за собой вращение валов $y_1, \dots, y_n$ в соответствии с уравнениями при любых данных начальных условиях. Будем говорить, что решение системы уравнений в полных дифференциалах возможно, если можно найти такую систему соединений, что любое вращение валов независимых переменных $x_1, \dots, x_m$ влечет за собой вращение валов зависимых переменных $y_1, \dots, y_n$ в соответствии с уравнениями при любых заданных начальных условиях. В большинстве теорем будут рассматриваться устройства, содержащие только интеграторы и сумматоры.

Основное условие разрешимости

Теорема 1. Для того чтобы систему обыкновенных дифференциальных уравнений можно было решить с применением только интеграторов и сумматоров, необходимо и достаточно, чтобы эти уравнения могли быть записаны в виде

$$\frac{dy_k}{dx_1} = \sum_{i,j=0}^n a_{ijk} y_i \frac{dy_j}{dy_1} \quad (k=2,3,\dots,n), \quad (1)$$

где $y_0 = 1$ (введено для компактности записи), y_1 — независимое переменное, $y_2, \dots, y_n$ — зависимые переменные, в число которых входят зависимые переменные первоначальной системы.

Доказательство. Условие (1) является необходимым. Действительно, предположим, что исходная система может быть решена с применением только интеграторов и сумматоров. Зависимые переменные должны появляться как выходы либо сумматоров, либо интеграторов. Можно считать, что все они являются выходами интеграторов. Для этого выходы сумматоров делаем переменными интегрирования на интеграторах, интегрируемые функции которых являются единичными константами. Итак, пусть имеется $n-1$ интеграторов. Обозначим их выходы через $y_2, \dots, y_n$. Каждое смещение (интегрируемая величина) должно вызываться одним из трех возможных источников: независимое переменное y_1 , выход интегратора или выход сумматора (в обобщенном смысле). Сумматор должен управляться совокупностью определенных функций y , включая, быть может, y_1 . Этим исчерпываются все частные случаи источника движения вида $\sum_{i=1}^N b_{ik} y_i$ для

k -го интегратора, где b_{ik} — константы. Очевидно, что b_{ik} могут представлять передаточные числа или сложные соединения сумматоров.

Без потери общности можно принять их равными 0 или 1. k -й интегратор может обладать в дополнение к значению источника движения еще и начальным смещением. Обозначим его через b_{0k} . Тогда интегрируемая функция k -го интегратора будет равна $\sum_{i=0}^n b_{ik} y_i$, где для удобства вводится $y_0 = 1$. По точно таким же соображениям переменные интегрирования на интеграторах имеют вид $\sum_{i=1}^n c_{jk} y_j$. Интеграторы накладывают на систему следующие требования:

$$y_k = \int \sum_i b_{ik} y_i d \sum_j c_{jk} y_j = \quad (2)$$

$$= \int \sum_i b_{ik} y_i \sum_j c_{jk} dy_j = \quad (3)$$

$$= \int \sum_{i,j} a_{ijk} y_i dy_j \quad (k = 2, 3, \dots, n), \quad (4)$$

где $a_{ijk} = b_{ik} c_{jk}$. Дифференцируя обе части по y_1 , получаем

$$\frac{dy_k}{dy_1} = \sum_{i,j} a_{ijk} y_i \frac{dy_j}{dy_1} \quad k = 2, 3, \dots, n.$$

Это и есть уравнения (1).

Докажем, что условие является и достаточным. Это следует из того, что система (1) может быть преобразована интегрированием к виду (4). Этот вид определяет устройство, использующее только интеграторы и сумматоры, причем каждый вал обладает не более чем одним источником движения.

Реализация функций

Будем говорить, что функция одного переменного $y = f(x)$ может быть реализована, если существует устройство, использующее только интеграторы и сумматоры и такое, что один вал может независимо вращаться как x , а другой вал вынужден вращаться как y . Из теоремы 1 следует, что если $f(x)$ может быть реализована, то должна существовать такая система уравнений (1), что если $y_1 = x$, то (скажем) y_2 равняется $f(x)$. Функция от n переменных $F(x_1, \dots, x_n)$ может быть реализована, если существует такое устройство, что n валов $x_1, \dots, x_n$ могут вращаться независимо, и еще один вал вращается как F .

Функции одного переменного классифицируются следующим образом. Если выполнено соотношение вида

$$y = a_0 + a_1 x + a_2 x^2 + \dots + a_m x^m, \quad (5)$$

где m — целое положительное число и $a_1, \dots, a_m$ — действительные константы, то y называется многочленом от x или целой рациональной функцией от x . Если

$$y = F_1(x)/F_2(x), \quad (6)$$

где F_1 и F_2 — многочлены от x , то y называется рациональной функцией от x . Если

$$R_0(x) + y R_1(x) + y^2 R_2(x) + \dots + y^n R_n(x) = 0, \quad (7)$$

где $R_0, \dots, R_n$ — рациональные функции от x , то y называется алгебраической функцией от x . Если не имеет места ни одно из соотношений такого вида, то y называется трансцендентной функцией от x . Трансцендентные функции могут быть разделены на два класса. Если удовлетворяется дифференциальное уравнение

$$\sum_j A_j x^{n_j} y^{n_{0j}} (y')^{n_{1j}} (y'')^{n_{2j}} \dots (y^{(m)})^{n_{mj}} = 0, \quad (8)$$

где $y^{(k)} = d^k y / dx^k$, A_j — константы и степени $n_j, n_{0j}, n_{1j}, \dots, n_{mj}$ — целые (т. е. если некоторый многочлен от $x, y, y', \dots, y^{(m)}$ тождественно равен нулю), то y называется алгебраической трансцендентной функцией от x . Если не существует ни одно из соотношений такого вида, то функция называется гипертрансцендентной, или трансцендентно-трансцендентной.

Очевидно, все алгебраические функции не являются гипертрансцендентными, так как посредством умножения (7) на все знаменатели и последующего дифференцирования получается выражение вида (8). В действительности лишь немногие из обычно встречающихся аналитических функций гипертрансцендентны; из них наиболее известны гамма-функция и дзета-функция Римана

$$\Gamma(n) = \int_0^{\infty} x^{n-1} e^{-x} dx; \quad \zeta(s) = \sum_{k=1}^{\infty} \frac{1}{k^s}.$$

Гипертрансцендентность первой функции была доказана Гольденом, а второй — Гильбертом. В табл. 1 приведена классификация некоторых из обычно встречающихся функций.

Таблица 1

Функции от одного переменного

Трансцендентные		Алгебраические	
Гипертрансцендентные	Алгебраические трансцендентные	Иррациональные алгебраические	Рациональные
Гамма-функция	Показательные и логарифмические	x^m , где m — рациональная дробь.	Частные от деления
Дзета-функция	Тригонометрические, гиперболические и обратные им. Функции Бесселя. Эллиптические функции и интегралы. Функция вероятности.	Решения алгебраического уравнения, выраженные через параметр.	Целые α, x, x^2 и др. многочлены.

Теорема 2. *Функция от одного переменного может быть реализована тогда и только тогда, когда она не является гипертрансцендентной.*

Доказательство. Покажем сначала, что если функция может быть реализована, то она не гипертрансцендентная. Любая функция $f(x)$, которая может быть реализована, должна удовлетворять системе уравнений вида (1) при $y_1 = x$ и $y_2 = f(x)$.

Дифференцируя систему (1) $n-2$ раз, получаем в итоге $(n-1)^2$ уравнений, из которых можно исключить $n^2 - 2n$ переменных

$$y_3, y'_3, \dots, y_3^{(n-1)}; y_4, y'_4, \dots, y_4^{(n-1)}, \dots; y_n, y'_n, \dots, y_n^{(n-1)}$$

с помощью, например, метода Сильвестра, требующего только умножения и сложения и поэтому дающего в результате соотношение вида (8) при $x = y_1$, $y = y_2$ и $1 = y_0$.

Для доказательства того, что любая негипертрансцендентная функция может быть реализована, покажем, что соотношение (8) может быть записано в виде (1). Пусть левая часть уравнения (8) равна Φ . Дифференцируя обе части уравнения (8) по x , получаем

$$\frac{\partial \Phi}{\partial x} + \frac{\partial \Phi}{\partial y} y' + \frac{\partial \Phi}{\partial y'} y'' + \dots + \frac{\partial \Phi}{\partial y^{(m)}} y^{(m+1)} = 0.$$

Всюду, за исключением точек, в которых $\frac{\partial \Phi}{\partial y^{(m)}} = 0$, имеем

$$y^{(m+1)} = \frac{\frac{\partial \Phi}{\partial x} + \frac{\partial \Phi}{\partial y} y' + \dots + \frac{\partial \Phi}{\partial y^{(m-1)}} y^{(m)}}{\frac{\partial \Phi}{\partial y^{(m)}}} = \frac{P_1(x, y, y', \dots, y^{(m)})}{P_2(x, y, y', \dots, y^{(m)})},$$

где P_1 и P_2 — многочлены от $x, y, y', y'', \dots, y^{(m)}$. Пусть $y_1 = x$, $y_2 = y$, $y_3 = y'$, $\dots$, $y_{m+3} = y^{(m+1)}$. Тогда получаем

$$\frac{dy_k}{dy_1} = y_{k+1}, \quad k = 2, 3, \dots, m+2$$

при дополнительном условии

$$y_{m+3} = \frac{P_1(y_1, y_2, \dots, y_{m+2})}{P_2(y_1, y_2, \dots, y_{m+2})}. \quad (9)$$

Теперь задача состоит в том, чтобы свести это соотношение к виду системы уравнений (1). Сначала рассмотрим функцию P_1 . Она является суммой произведений переменных $y_1, \dots, y_{m+3}$ (так как степени целые, то они могут рассматриваться как произведения повторяющихся множителей). Пусть первый член числителя равен $u_1 u_2 \dots u_s$, где $u_1, u_2, \dots, u_s$ — это некоторые из переменных $y_1, \dots, y_{m+3}$, и пусть

$$\frac{dy_{m+4}}{dy_1} = u_1 \frac{du_2}{dy_1} + u_2 \frac{du_1}{dy_1},$$

так что $y_{m+4} = u_1 u_2$. Далее пусть

$$\frac{dy_{m+5}}{dy_1} = u_3 \frac{dy_{m+4}}{dy_1} + y_{m+4} \frac{du_3}{dy_1}.$$

Следовательно, $y_{m+5} = y_{m+4} u_3 = u_1 u_2 u_3$. Продолжая таким образом, получаем, наконец,

$$y_{m+s+2} = u_1 u_2 \dots u_s.$$

Точно так же поступаем с каждым членом числителя и знаменателя, продолжая ряд уравнений до тех пор, пока каждому произведению из P_1 и P_2 не будет поставлено в соответствие некоторое переменное y . Обозначим переменные y , соответствующие членам P_1 , через $v_1, v_2, \dots, v_r$, а соответствующие членам P_2 — через $w_1, w_2, \dots, w_t$. Тогда уравнение (9) сводится к условию

$$y_{m+3} = \sum_{h=1}^r v_h \left/ \sum_{h=1}^t w_h \right.$$

Заключительный шаг доказательства состоит в сведении этой системы к виду (1). Предположим, что последнее из переменных y , т. е. w_t , есть y_{g-1} . Пусть

$$\begin{aligned} \frac{dy_g}{dy_1} &= -y_{g+1} \frac{d \sum w}{dy_1}, \\ \frac{dy_{g+1}}{dy_1} &= 2y_g \frac{dy_g}{dy_1}. \end{aligned}$$

Следовательно, $y_g = 1/\sum w$. Далее, положив

$$\frac{dy_{m+3}}{dy_1} = y_g \frac{d \sum v}{dy_1} + \sum v \frac{dy_g}{dy_1},$$

получаем $y_{m+3} = \sum v / \sum w$. Тем самым уравнение (8) сведено к виду (1).

Теорема 3. Если функция одного переменного $y = f(x)$ может быть реализована, то могут быть реализованы ее производная $z = f'(x)$, ее интеграл $w = \int_{\alpha}^x f(x) dx$ и обратная ей функция $x = f^{-1}(y)$.

Для доказательства первой части возьмем производную от (8). В результате получаем два уравнения, из которых можно исключить y посредством преобразований, включающих только умножение и сложение. Результатом является соотношение вида (8) без членов, содержащих y . Заменяя y' на z , y'' на z' и т. д., убеждаем-

ся, что функция z не является гипертрансцендентной, если y не гипертрансцендентная.

Для доказательства второй части просто заменим в соотношении (8) y на w' , y' на w'' и т. д. Таким образом, получаем другое соотношение того же вида. Следовательно, если функция y не гипертрансцендентная, то и функция w не гипертрансцендентная.

Для доказательства третьей части заменим y' на $1/x'$ (где $x' = dx/dy$) y'' на $-x''/(x')^3$ и т. д. Так как все производные от функции y могут быть выражены как отношения многочленов от производных от x , то получающееся в результате выражение может быть приведено посредством умножения на все знаменатели к виду (8) с взаимной заменой x и y . Это означает, что функция, обратная негипертрансцендентной, также не гипертрансцендентная.

Теорема 4. *Если могут быть реализованы функции f и g , то может быть реализована и их суперпозиция $y = f(g(x))$.*

Доказательство. Утверждение теоремы будет доказано таким образом: две системы уравнений, каждая из которых имеет вид (1), будут записаны как единая система того же вида, включающая большее число уравнений. Предположим, что функция g удовлетворяет системе данного типа при $y_1 = x$ и $y_2 = g$, причем k пробегает значения от 2 до n . Функция f также удовлетворяет системе этого типа. Так как аргументом функции f является g , то заменим y_1 в этой системе на y_2 . Если индексы при y в системе для f пробегает значения от $n+1$ до $n+m$ (причем m — число уравнений в системе для f), то получим систему уравнений вида (1) при $y_1 = x$ и $y_{m+2} = f(g(x))$.

Хотя, как вытекает из теоремы 2, могут быть точно реализованы лишь функции, не являющиеся гипертрансцендентными, можно аппроксимировать значительно более широкий класс функций, используя только интеграторы.

Теорема 5. *Любая функция $f(x)$, непрерывная в замкнутом интервале $a \leq x \leq b$, может быть реализована в этом интервале с точностью до заданного значения допустимой ошибки $\varepsilon > 0$ с применением только конечного числа интеграторов. Это означает, что может быть найдено устройство, реализующее такую функцию $F(x)$, что*

$$|F(x) - f(x)| < \varepsilon,$$

при $a \leq x \leq b$.

Доказательство. Наше доказательство основывается на знаменитой теореме Вейерштрасса, которая утверждает, что

любую функцию $f(x)$ такого типа можно аппроксимировать многочленом $F(x)$ степени n

$$F(x) = \sum_{k=0}^n a_k x^k \quad (10)$$

при достаточно большом значении n .

Пусть теперь

$$dy_{j+1}/dy_j = a_j j! + y_{j+2}, \quad j = 2, 3, \dots, n; \quad y_{n+2} = 0.$$

Эта система, имеющая вид (1), удовлетворяет соотношению (10), если положить $y_1 = x$ и $y_2 = F(x)$. Кроме того, для устройства, реализующего решение, требуются только интеграторы, причем аддитивные константы определяются просто их начальными установками. Следовательно, теорема верна.

Если допустить останов машины и вращение валов вручную, то можно очевидным образом расширить условия теоремы 5 на все функции, непрерывные всюду, кроме конечного числа точек конечного разрыва.

Теперь перейдем к обобщению некоторых из этих понятий и теорем на случай функций, зависящих более чем от одного переменного.

Теорема 6. *Функция $y_{m+1} = f(y_1, \dots, y_m)$ от m переменных может быть реализована тогда и только тогда, когда она удовлетворяет системе уравнений в полных дифференциалах вида*

$$dy_k = \sum_{i,j=0}^n a_{ijk} y_i dy_j \quad k = m+1, m+2, \dots, n, \quad (11)$$

где $y_0 = 1$ и все a_{ijk} — действительные константы.

Доказательство необходимости и достаточности проводится в точности по тому же плану, который был применен при доказательстве теоремы 1. Было показано, что решение уравнений (1) является не гипертрансцендентной функцией от одного переменного. Теперь (11) может рассматриваться как обобщение (1). Будем говорить, что функция от m переменных, удовлетворяющая системе уравнений (11), является не гипертрансцендентной функцией от этих m переменных. При этом определении получаем в качестве обобщения теоремы 2 утверждение, что функция от m переменных может быть реализована тогда и только тогда, когда она не является гипертрансцендентной функцией от этих переменных. Очевидно, что для того, чтобы функция от m переменных не была гипертрансцендентной, необходимо, чтобы она не была гипертрансцендентной для каждого отдельного переменного, когда все остальные переменные заменены произвольными константами.

Таким образом, $x + \Gamma(y)$ является гипертрансцендентной функцией от x и y , так как замена x на 0 дает гипертрансцендентную функцию от одного переменного $\Gamma(y)$. Примерами негипертрансцендентных функций более чем от одного переменного являются $x + y$, $x \cdot y$, x^y , $\log_x y$, а также суперпозиции этих функций и негипертрансцендентных функций от одного переменного.

Обобщением теоремы 4 является следующее утверждение.

Теорема 7. Если две функции $f(x_1, \dots, x_n)$ и $g(y_1, \dots, y_m)$ от нескольких переменных могут быть реализованы, то можно также реализовать любую суперпозицию этих функций, например $\Phi(x_2, x_3, \dots, x_n, y_1, y_2, \dots, y_m) = f(g(y_1, \dots, y_m), x_2, \dots, x_n)$.

Это утверждение может быть доказано тем же методом, что и теорема 4, т. е. путем объединения двух систем уравнений в одну систему типа (11), содержащую большее число уравнений. Теорема 5 может быть обобщена на случай функций более чем от одного переменного, но при этом уже недостаточно одних интеграторов.

Теорема 8. Если задана любая функция $f(x_1, \dots, x_n)$ от n переменных, непрерывная по всем аргументам в замкнутой области n -мерного пространства $a_k \leq x_k \leq b_k$, $k = 1, 2, \dots, n$, то можно реализовать с применением только конечного числа интеграторов и сумматоров такую функцию $F(x_1, \dots, x_n)$, что в области $a_k \leq x_k \leq b_k$

$$|f - F| < \varepsilon,$$

где ε — заданное сколь угодно малое положительное число.

Доказательство. В силу обобщения теоремы Вейерштрасса функцию f можно описанным способом аппроксимировать многочленом F от переменных $x_1, \dots, x_n$. Так как, согласно нашему определению, многочлен не является гипертрансцендентной функцией, то он может быть реализован при помощи интеграторов и сумматоров. Следовательно, теорема доказана.

Первая часть теоремы 3 может быть обобщена на функции более чем от одного переменного следующим образом.

Теорема 9. Если функция $f(x_1, \dots, x_n)$ от n переменных может быть реализована, то может быть реализована и ее частная производная по любому переменному, например по x_1 .

Доказательство. Из теоремы 6 следует, что если функция может быть реализована, то она удовлетворяет системе уравнений вида

$$dy_i = \sum_{j=1}^m A_{ij} dx_j + \sum_{j=1}^s B_{ij} dy_j, \quad i = 1, 2, \dots, s,$$

где $x_1, \dots, x_m$ — независимые переменные, $y_1, \dots, y_s$ — зависимые переменные, A_{ij} и B_{ij} — линейные формы от этих переменных и, скажем, y_1 равняется f . Разделив эти уравнения на dx_1 и положив $0 = dx_2 = dx_3 = \dots = dx_m$, получаем

$$\frac{\partial y_i}{\partial x_1} = A_{i1} + \sum_{j=1}^s B_{ij} \frac{\partial y_j}{\partial x_1}, \quad i = 1, 2, \dots, s,$$

или

$$\sum_{j=1}^s (B_{ij} - \delta_{ij}) \frac{\partial y_j}{\partial x_1} = A_{i1}, \quad i = 1, 2, \dots, s.$$

Значение $\frac{\partial y_1}{\partial x_1}$ по правилу Крамера равно

$$\frac{\partial y_1}{\partial x_1} = \frac{\begin{vmatrix} A_{11} & B_{12} & \dots & B_{1s} \\ A_{21} & B_{22} & \dots & B_{2s} \\ \dots & \dots & \dots & \dots \\ A_{s1} & B_{s2} & \dots & B_{ss} \end{vmatrix}}{|B_{ij} - \delta_{ij}|} = \frac{P_1}{P_2},$$

где P_1 и P_2 — многочлены от данных переменных.

Это уравнение может быть сведено к виду системы уравнений (11) в точности тем же самым методом, который был использован при сведении уравнения (9), которое имело такой же вид.

Последняя часть теоремы 3 также может быть обобщена следующим образом.

Теорема 10. Если может быть реализована функция $y = f(x_1, \dots, x_n)$ от n переменных, то может быть реализована и функция, обратная данной относительно любого одного переменного, например, $x_1 = F(y, x_2, \dots, x_n)$.

Доказательство. Взяв полный дифференциал от y , получаем

$$dy = \sum_{i=1}^s f'_{x_i} dx_i.$$

Следовательно,

$$dx_1 = \frac{1}{f'_{x_1}} \left(dy - \sum_{i=2}^s f'_{x_i} dx_i \right),$$

т. е.

$$x_1 = \int \frac{1}{f'_{x_1}} \left(dy - \sum_{i=2}^s \frac{f'_{x_i}}{f'_{x_1}} dx_i \right).$$

Так как функция f может быть реализована, то в силу предыдущей теоремы могут быть реализованы члены $f'_{x_i} (i = 1, 2, \dots, n)$. Обратные дроби и частные не являются гипертрансцендентными функциями, а, следовательно, члены $1/f_{x_1}$ и $-f'_{x_i}/f_{x_1}$ могут быть реализованы. Отсюда следует, что величина x_1 может быть получена посредством реализации этих подинтегральных выражений, интегрирования их по соответствующим переменным и сложения результатов.

Системы уравнений

Теперь можно доказать следующую общую теорему о дифференциальном анализаторе.

Теорема 11. Самая общая система обыкновенных дифференциальных уравнений

$$f_k(x; y_1, y'_1, \dots, y_1^{(m)}; y_2, y'_2, \dots, y_2^{(m)}; \dots; y_n, y'_n, \dots, y_n^{(m)}) = 0$$

$$k = 1, 2, \dots, n \quad (12)$$

порядка m с n зависимыми переменными может быть решена на дифференциальном анализаторе с использованием только конечного числа интеграторов и сумматоров при условии, что функции f_k являются суперпозициями негипертрансцендентных функций от данных переменных.

Перед тем как доказывать эту теорему, необходимо вывести предварительную лемму. Естественный метод состоял бы в том, чтобы разрешить уравнение $f_1 = 0$ относительно $y_1^{(m)}$, уравнение $f_2 = 0$ относительно $y_2^{(m)}$ и т. д.

$$y_k^{(m)} = \Phi_k(x; y_1, y'_1, \dots, y_1^{(m)}; y_2, y'_2, \dots, y_2^{(m)}; \dots; y_n, y'_n, \dots, y_n^{(m)}).$$

Однако может оказаться, что $y_1^{(m)}$ не входит в f_1 , но входит в некоторую другую функцию. Покажем сначала, что, продифференцировав уравнения (12) и изменив порядок уравнений, можно получить эквивалентную систему, при которой наивысший порядок производной от y_1 встречается в первом уравнении, наивысший порядок производной от y_2 встречается во втором уравнении и т. д.

Прежде всего заметим, что если f_k рассматривается как функция от независимых переменных

$$(x, y_1, y'_1, \dots, y_1^{(m)}; y_2, y'_2, \dots, y_2^{(m)}; \dots; y_n, y'_n, \dots, y_n^{(m)}),$$

то дифференцирование уравнения $f_k = 0$ дает эквивалентное уравнение (при условии, что граничные условия выбираются в соответствии с исходным уравнением). В этом уравнении наивысший порядок производных от каждого встречающегося переменного увеличен на единицу. Кроме того, если исходная функция f_k не включала

гипертрансцендентных функций, то по теореме 9 полученная функция также не включает их и может быть реализована.

Для наших целей существенной характеристикой уравнения (12) является набор значений высших порядков производных различных переменных, встречающихся в этих уравнениях. Эти значения могут быть заданы квадратной матрицей следующим образом:

$$\begin{array}{c|cccc}
 & y_1 & y_2 & \dots & y_n \\
 \hline
 f_1 & a_{11} & a_{12} & \dots & a_{1n} \\
 f_2 & a_{21} & a_{22} & \dots & a_{2n} \\
 f_3 & a_{31} & a_{32} & \dots & a_{3n} \\
 \vdots & & \dots & & \dots \\
 \vdots & & \dots & & \dots \\
 f_n & a_{n1} & a_{n2} & \dots & a_{nn}
 \end{array}$$

Здесь a_{jk} — самый высокий порядок производной y_k , встречающийся в функции f_j . Наши a являются целыми числами, которые могут принимать значения от 0 до m . Если же переменное вовсе не встречается, то можно придать a специальный символ λ . Дифференцирование функций f_k дает в результате прибавление единицы к каждому элементу k -й строки, кроме элементов λ , которые остаются неизменными. Две строки матрицы могут быть переставлены, так как это просто означает перенумерацию функций. Покажем, что путем переупорядочивания и дифференцирования всегда можно найти новую систему, причем a_{11} является максимальным числом в первом столбце, a_{22} — максимальным числом во втором столбце и в общем случае a_{kk} — максимальным числом в k -м столбце.

Это будет показано методом математической индукции. Покажем сначала, что это верно при $n = 2$. Для двух переменных имеем матрицу

$$\begin{array}{c|cc}
 & y_1 & y_2 \\
 \hline
 f_1 & a_{11} & a_{12} \\
 f_2 & a_{21} & a_{22}
 \end{array}$$

Если λ встречается больше одного раза, то система вырожденная. Если одна из букв, например a_{11} , есть λ , то можно поменять местами строки и затем дифференцировать вторую строку до получения для нового значения a_{22} неравенства $a_{22} \geq a_{12}$. Если ни одна из букв не есть λ , то либо $a_{22} = a_{12}$, либо одно из этих значений меньше. Если они равны, то строки могут быть переставлены

в случае необходимости для получения неравенства $a_{11} \geq a_{21}$. Если же одно из них меньше, то будем дифференцировать соответствующую строку до тех пор, пока они станут равны, а затем поступим, как раньше. Таким образом, лемма верна при $n = 2$.

Теперь, предположив, что утверждение верно для n , покажем, что оно верно для $n + 1$, и тем самым завершим доказательство. Согласно индуктивному предположению, если дана любая квадратная матрица, содержащая $(n + 1)^2$ элементов, можно путем дифференцирования соответствующей системы уравнений и перепорядочения первых n строк найти эквивалентную систему

	y_1	y_2	$\dots$	y_n	y_{n+1}
f_1	a_{11}	a_{12}	$\dots$	a_{1n}	$a_{1, n+1}$
f_2	a_{21}	a_{22}	$\dots$	a_{2n}	$a_{2, n+1}$
$\vdots$	$\dots$	$\dots$	$\dots$	$\dots$	$\dots$
f_n	a_{n1}	a_{n2}	$\dots$	a_{nn}	$a_{n, n+1}$
f_{n+1}	$a_{n+1,1}$	$a_{n+1,2}$	$\dots$	$a_{n+1,n}$	$a_{n+1, n+1}$

такую, что $a_{jk} \geq a_{ik}$ при $k, j = 1, 2, \dots, n$. Можно также предположить, что $a_{kk} > a_{n+1,k}$ для любого $k \leq n$, так как в противном случае можно добиться этого путем одновременного дифференцирования всех первых n функций. Теперь возникают две возможности: 1) $a_{n+1, n+1} \geq a_{j, n+1}$; $j = 1, 2, \dots, n$. В этом случае система в имеющемся виде является удовлетворительной. 2) Это условие не имеет места; тогда существует некоторое значение $a_{j, n+1} > a_{n+1, n+1}$. Предположим, что это $a_{1, n+1}$. Будем теперь дифференцировать последнюю строку до тех пор, пока не получим выполнение одной из следующих трех возможностей: 1) $a_{n+1, n+1} = a_{1, n+1}$ (в этом случае система является удовлетворительной); 2) $a_{n+1, 1} = a_{11}$ (в этом случае меняем местами первую и последнюю строки, после чего система является удовлетворительной); 3) $a_{n+1, s} = a_{ss}$ для одного или нескольких значений s между 1 и $n + 1$. В этом случае дифференцируем оба уравнения с номерами s и $n + 1$ до тех пор, пока не произойдет одно из трех: а) $a_{n+1, n+1} = a_{1, n+1}$ или $a_{n+1, 1} = a_{11}$, тогда действуем как выше при 1) и 2); б) $a_{s1} = a_{11}$ или $a_{s, n+1} = a_{1, n+1}$, тогда меняем местами s -ю и $(n + 1)$ -ю строки и поступаем, как в случае 1); в) максимальное значение в некотором другом столбце достигнуто соответствующими элементами в дифференцируемых строках. В этом случае присоединяем строку или строки, содержащие это максимальное значение, к тем, которые уже дифференцировались, и продолжаем действовать тем же методом. Легко видеть, что этот процесс должен сойтись после конеч-

ного числа шагов, так как невозможно, чтобы все элементы первого и последнего столбцов, кроме a_{11} и $a_{1,n+1}$, были λ . Этим завершается доказательство.

Теперь сравнительно просто доказать теорему 11. Прежде всего найдем только что описанным методом систему, эквивалентную (12), в которой самый высокий порядок производной y_k появляется в уравнении с номером k . Пусть этот порядок равен p_k . Было показано, что может быть реализована любая суперпозиция негипер-трансцендентных функций (теоремы 4 и 7), а следовательно, могут быть реализованы функции f_k . Следовательно, по теореме об обратных функциях (теорема 10), можно реализовать функции

$$y_k^{(p_k)} = \Psi_k(x, y_1, y_1', \dots, y_2, y_2', \dots, y_n, y_n', \dots).$$

Переменные $y_k^{(v)}$ при $v < p_k$ могут быть получены интегрированием $y_k^{(p_k)}$ и, следовательно, рассматриваемая система может быть решена с применением только интеграторов и сумматоров.

Аппроксимация множителей с помощью передаточных чисел

При подготовке дифференциального анализатора к решению задач часто бывает необходимо вводить постоянные множители. Это можно делать различными способами. Как указывалось, интегратор с заданной постоянной интегрируемой величиной k выдает

$$\omega = \int_0^v k dv = kv. \text{ Второй метод состоит в том, чтобы получить}$$

грубое приближение k с помощью шестерен с общим передаточным числом, например k' . Другой набор шестерен используется для получения грубого приближения k'' к значению $(k - k')$. Переменная $k'x$ объединяется с $k''x$ через сумматор для получения приближения второго порядка к k . При единственном предположении, что имеется неограниченное число сумматоров и пар шестерен одного размера с передаточным числом $a \neq 0, 1$, легко показать, что, действуя таким образом, можно получить такое передаточное число t , что $|t - k| < \varepsilon$, где ε — сколь угодно малое заданное положительное число. Третий способ получения передаточного числа k состоит в использовании только шестерен. Покажем, что любое передаточное число может быть аппроксимировано с применением только конечного числа пар шестерен двух размеров. Точнее говоря, справедлива следующая теорема.

Теорема 12. Если даны два передаточных числа a и b пар шестерен, ни одно из которых не равно ни нулю, ни единице, и если b не является рациональной степенью числа a , то можно найти

такие положительные или отрицательные целые числа u и v , что $0 \leq |a^u b^v - k| < \varepsilon$, где ε — заданное произвольно малое число.

Без ограничения общности полагаем $a > 1$; в противном случае изменение порядка шестерен в паре обеспечивает передаточное число $1/a > 1$. Сначала покажем, что может быть получено такое передаточное число U , что $1 < U < 1 + \delta$ при сколь угодно малом значении δ . Для этого должны быть найдены такие целые числа x и y , что $1 < a^x b^y < 1 + \delta$. Так как $a > 1$, то можно произвести логарифмирование по основанию a , сохраняя тот же порядок неравенств:

$$0 < x + y \log_a b < \log_a (1 + \delta) = \mu.$$

Так как $\log_a b$ — иррациональное число, то можно удовлетворить этому неравенству в силу известных теорем о диофантовом приближении. Таким образом, можно получить значение U передаточного числа шестерен. Если наборы шестерен с передаточным числом U последовательно связаны друг с другом r раз, причем r выбрано так, что

$$U^{r-1} < k \leq U^r,$$

то разность между U^r и k меньше, чем δk , а следовательно, сколь угодно мала.

Если b — рациональная степень a , например $a^{m/n} = b$, где m/n — несократимая дробь, то необходимое и достаточное условие возможности получить какое-либо передаточное число состоит в том, чтобы это число имело вид $a^{k/n}$, где k — любое целое число. Прежде всего, любое передаточное число выражается как $a^x b^y = a^{(xn+ym)/n}$, т. е. имеет указанный вид. Достаточность условия следует из того факта, что диофантово уравнение $xn + ym = k$ имеет решение в целых числах при любом целом k , если n и m взаимно просты.

Автоматическое управление скоростью работы устройства

Важной частью управляющей схемы дифференциального анализатора является автоматическое управление скоростью его работы. Выходы интеграторов обладают максимальной скоростью s , при превышении которой вероятно появление проскальзывания или вибрации. Для того чтобы воспрепятствовать превышению этой скорости, следовало бы вращать вал независимого переменного с такой скоростью, чтобы предельная скорость вращения валов интеграторов достигалась только в случае максимального смещения. Однако при такой системе решение заняло бы больше

времени, чем необходимо, так как на протяжении большинства решений ни один из интеграторов не находился бы в максимальном смещении. Система автоматического управления скоростью является устройством, заставляющим вал независимого переменного вращаться с такой скоростью, что интегратор, вращающийся быстрее всех, движется со скоростью s . При использовании некоторых довольно широких предположений можно получить приближенную оценку времени, которое экономится этим способом. Результаты приводимой ниже теоремы сравнивались с несколькими экспериментальными решениями, и ошибка во всех случаях была меньше 7%.

Теорема 13. Пусть n интеграторов приводятся в действие независимым переменным с максимальной скоростью s и интегрирующим множителем b ; предполагается, что смещения интеграторов случайно распределены по отношению к независимому переменному в пределах от $-a$ до $+a$ и скорость ограничена только этими пределами и независимым переменным x с максимальной скоростью r . Тогда ожидаемая средняя скорость независимого переменного равна

$$\bar{\dot{x}} = \frac{1}{[nab/x(n+1)s] + [1/r(n+1)] [(s/abr)^n]}$$

Доказательство. Возьмем

$$\dot{x} = \frac{\int\limits_{x=0}^{x=u} dx}{\int\limits_{x=0}^{x=u} dt}.$$

Здесь u — максимальное значение x , t — время. Далее, так как $dt = \frac{dx}{dx/dt}$, имеем

$$\bar{\dot{x}} = \frac{u}{\int\limits_0^u \frac{dt}{dx} dx}.$$

Пусть y — смещение интегратора с наибольшим смещением. Тогда вероятность того, что оно заключено между y и $y+dy$, равна

$$n(y/a)^{n-1} dy/a.$$

Ожидаемое число оборотов вала x при этом условии равно

$$dx = n(y/a)^{n-1} u dy/a.$$

Если $y \leq s/br$, то скорость ограничена независимым переменным и dx/dt равняется r . Если же $y > s/br$, то скорость ограничена интеграторами и $dx/dt = x/yb$. В первом случае

$$dt = un/r (y/a)^{n-1} dy/a.$$

Во втором случае

$$dt = yb/snu (y/a)^{n-1} dy/a.$$

Следовательно,

$$\dot{x} = \frac{u}{\int_{y=0}^{y=s/br} \frac{un}{ar} \left(\frac{y}{a}\right)^{n-1} dy + \int_{y=s/br}^{y=a} \frac{bnu}{s} \left(\frac{y}{a}\right)^n dy},$$

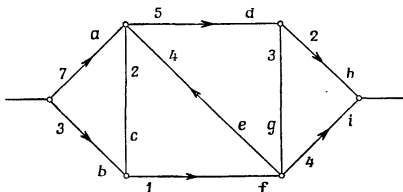
или

$$\dot{x} = \frac{1}{\frac{nab}{(n+1)s} + \frac{1}{r(n+1)} \left(\frac{s}{abr}\right)^n}.$$

С л е д с т в и е. Если максимальную скорость независимого переменного сделать бесконечной, то время решения уменьшится в $\frac{n+1}{n}$ раз по сравнению с тем значением, которое получилось бы в случае, если бы валу x была придана такая постоянная скорость $v = s/ab$, что при максимальном смещении интеграторы как раз достигали бы максимальной допустимой скорости.

Этот факт следует из общего выражения при r , стремящемся к бесконечности.

Рассмотрим некоторую двухполюсную сеть, например изображенную на рис. 1. Ребра такой сети могут представлять каналы связи или вообще любые проводящие системы ограниченной пропускной способности, как, например, железнодорожные магистрали, системы энергоснабжения или трубопроводы, при условии, что в каждом случае возможно установить определенную максимально-



Р и с. 1.

допустимую величину потока через данное ребро — пропускную способность ребра. Ребра могут быть двух типов: ориентированные, направление которых указывается стрелкой, и неориентированные, причем в последнем случае в любом направлении допустим поток произвольной величины вплоть до пропускной способности. В вершинах, или точках соединения сети, допускается любое перераспределение входящего потока в выходящий поток, подчиненное лишь тому ограничению, что в каждом ребре не превышаетея пропускная способность и в каждой вершине выполняется закон Кирхгофа, согласно которому суммарный (алгебраический) поток через вершину равен нулю. Заметим, что в случае информационного потока это может потребовать сколь угодно больших задержек в каждой

¹⁾ Elias P., Feinstein A. and Shannon C., A note on the maximum flow through a network, *Trans. IRE*, IT-2, 4 (1956), 117.

вершине, чтобы допустить перекодирование сигналов, выходящих из этой вершины. Задача состоит в том, чтобы вычислить максимальный поток через сеть в целом, входящий в левом полюсе и выходящий в правом полюсе сети.

Ответ можно получить в терминах сечений сети. *Сечение* двухполюсной сети есть совокупность ребер, при удалении которых сеть распадается на две или большее число несвязанных частей с двумя полюсами в разных частях. Таким образом, каждый путь от одного полюса к другому в исходной сети проходит хотя бы через одно ребро сечения. В приведенной выше сети примерами сечений являются (d, e, f) , (b, c, e, g, h) , (d, g, h, i) . Под *простым сечением* будем подразумевать такое сечение, что если удалить любое его ребро, то оно перестанет быть сечением. Так (d, e, f) и (b, c, e, g, h) являются простыми сечениями, в то время как (d, g, h, i) не является таковым. Если в связной двухполюсной сети удаляется простое сечение, то сеть распадается ровно на две части: *левую часть*, содержащую левый полюс, и *правую часть*, содержащую правый полюс. Каждому простому сечению припишем *значение*, равное сумме пропускных способностей ребер этого сечения, причем пропускная способность ориентированного ребра учитывается только тогда, когда оно направлено из левой части сети в правую часть. Заметим, что направление ориентированного ребра нельзя, вообще говоря, вывести по его положению в графе, соответствующем сети. Ребро направлено слева направо в простом сечении тогда и только тогда, когда стрелка на этом ребре направлена из вершины левой части сети к вершине правой части. Так, в нашем примере сечение (d, e, f) имеет значение $5 + 1 = 6$, а сечение (b, c, e, g, h) имеет значение $3 + 2 + 3 + 2 = 10$.

Теорема. *Максимальный поток слева направо через сеть равен минимальному значению всех простых сечений.*

Эта теорема может показаться почти очевидной из физических соображений и в течение некоторого времени она, по-видимому, принималась в области теории связи без доказательства. Однако, хотя первое утверждение, что этот поток не может быть превышен, в самом деле почти очевидно, второе утверждение, что он действительно может быть достигнут, никоим образом не является очевидным. Нам известно, что доказательства этой теоремы были даны Фордом и Фалкерсоном¹⁾ и Данцигом и Фалкерсоном²⁾. Ниже-

1) Ford L., Fulkerson D., Maximal flow through a network, *Can. J. Math.*, 51 (1956), 399.

2) Dantzig G., Fulkerson D., On the max flow min cut theorem of networks, Linear inequalities and related systems, *Ann. of Math. Study*, 38 (1956), 215. (Русский перевод: «Теорема о максимальном потоке и минимальном разрезе в сетях», сб. «Линейные неравенства», ИЛ, М., 1959.)

следующее доказательство относительно просто и, думается, существенно отличается от упомянутых.

Для того чтобы доказать сначала, что поток через сеть не может превышать значения минимального сечения¹⁾, рассмотрим произвольный заданный поток и сечение S с минимальным значением. Образует алгебраическую сумму S потоков слева направо через это сечение. Ясно, что она меньше или равна значению V этого сечения, так как последнее получилось бы, если бы все пути слева направо в S были загружены вплоть до пропускной способности, а пути в обратном направлении не использовались. Прибавим теперь к S алгебраическую сумму потоков во всех вершинах правой относительно сечения S части сети. Эта сумма равна нулю ввиду того, что закон Кирхгофа соблюдается в каждой вершине. С другой стороны, видно, что эта сумма уничтожает каждый поток, представленный в сумме S , а также, что каждый поток по ребру, оба конца которого лежат в правой части, встречается как со знаком плюс, так и со знаком минус и, следовательно, уничтожается. Единственная остающаяся величина, которая не уничтожается, есть поток из правого полюса сети или, как говорят, суммарный поток F через сеть. Заключаем отсюда, что $F \leq V$.

Докажем теперь более интересное положительное утверждение этой теоремы о том, что можно найти поток, который действительно достигает величины V . По любой заданной сети со значением V минимального сечения можно построить *приведенную* сеть, обладающую следующими свойствами:

- 1) граф приведенной сети есть граф исходной сети, за исключением, быть может, того, что несколько ребер исходной сети отсутствуют (нулевая пропускная способность) в приведенной сети;
- 2) каждое ребро в приведенной сети имеет пропускную способность, не большую, чем у соответствующего ребра исходной сети;
- 3) каждое ребро приведенной сети входит²⁾ по крайней мере в одно сечение со значением V , причем V есть минимальное значение сечения приведенной сети.

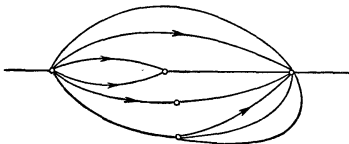
Приведенную сеть можно построить следующим образом. Если существует какое-нибудь ребро, которое не входит ни в какое минимальное сечение, то будем уменьшать его пропускную способность до тех пор, пока это ребро не войдет в какое-нибудь минимальное сечение или пока его пропускная способность не достигнет нуля. В последнем случае удалим это ребро из сети.

¹⁾ То есть простого сечения с минимальным значением. В дальнейшем авторы, говоря о простом сечении, опускают слово «простой». — *Прим. перев.*

²⁾ Здесь выражение «ребро входит в сечение» нужно понимать в том смысле, что ребро принадлежит сечению, имеет ненулевую пропускную способность и направлено слева направо в этом сечении, если оно является ориентированным. — *Прим. перев.*

Далее, возьмем любое другое ребро, не входящее ни в какое минимальное сечение, и выполним ту же самую операцию. Продолжим этот процесс до тех пор, пока не останется ребер, не входящих в какое-нибудь минимальное сечение. Ясно, что полученная сеть удовлетворяет требуемым условиям. Вообще говоря, имеется несколько различных приведенных сетей, которые можно получить из данной сети в зависимости от того порядка, в котором выбираются ребра. Если искомый поток можно найти в приведенной сети, то ясно, что тот же самый поток будет искомым и в исходной сети, потому что как условие Кирхгофа, так и условие ограничения пропускной способности будут выполнены. Следовательно, если доказать теорему для приведенных сетей, то она будет вообще справедлива.

Доказательство проведем индукцией по числу ребер. Заметим сначала, что если каждый путь через приведенную сеть содержит

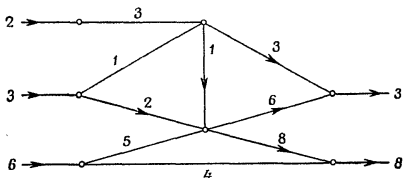


Р и с. 2.

только два или меньшее число ребер, то эта сеть имеет вид, указанный на рис. 2. Вообще, такая сеть состоит из параллельного соединения последовательных подсетей, причем эти подсети имеют пути, состоящие самое большее из двух ребер со стрелками слева направо или без них. Очевидно, что для приведенной сети такого вида теорема справедлива. Необходимо только загрузить каждое ребро вплоть до его пропускной способности. Предположим теперь, что теорема справедлива для всех приведенных сетей, имеющих менее n ребер. Покажем, что она справедлива для любой приведенной сети с n ребрами.

Возможны два случая: либо данная приведенная сеть с n ребрами имеет путь слева направо длины по крайней мере три, либо она является сетью только что описанного вида. В последнем случае, как было упомянуто выше, теорема справедлива. В первом случае, если взять второе ребро на пути длины три или более, то получим элемент, лежащий между внутренними вершинами. Существует (так как сеть приведенная) минимальное сечение S , содержащее это ребро. Заменим каждое ребро этого сечения двумя

последовательными ребрами, каждое из которых имеет ту же самую пропускную способность, что и исходное ребро. Затем соединим вместе (отождествим) все эти вновь образованные средние вершины в одну общую вершину. В результате этого сеть станет последовательным соединением двух более простых сетей. Каждая из них имеет то же самое значение V минимального сечения, поскольку и та и другая содержат сечение, соответствующее сечению C , и к тому же никакая из них не может содержать сечений с меньшими



Р и с. 3.

значениями, так как указанная операция отождествления вершин может только устранять сечения, но не может вводить новых сечений.

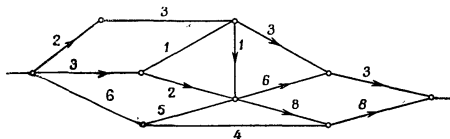
Каждая из этих двух последовательных сетей содержит менее n ребер. Это следует из наличия пути из левого полюса в правый длины по крайней мере три. Наличие этого пути влечет за собой существование ребра в левой части, которое не встречается в правой части, и наоборот. Таким образом, по индуктивному предположению можно создать искомый суммарный поток в каждой из этих сетей. Очевидно, что когда общая вершина разъединена и сеть возвращена к своему первоначальному виду, то тот же самый поток останется неизменным и в исходной сети. Этим завершается доказательство.

Интересно, что в приведенной сети каждое ребро загружается вплоть до своей пропускной способности и направление потока определяется любым минимальным сечением, содержащим это ребро. В неприведенных сетях, вообще говоря, имеется некоторая свобода в выборе потока по ребрам и даже иногда в выборе направления потока.

К вышеизложенному результату можно легко свести более общую задачу относительно потока через сеть. Предположим, что имеется сеть с несколькими входами и выходами, как показано на рис. 3. Три вершины слева являются входами, и желательно ввести через

эти входы 2, 3 и 6 единиц потока. Вершины справа — выходы, и желательно через эти вершины передать 3 и 8 единиц потока. Задача состоит в том, чтобы найти условия, при которых это возможно.

Эту задачу можно свести к предыдущей соединением входных каналов в общий левый узел и соединением выходных каналов в общий правый узел. В нашем частном случае это приводит к сети, изображенной на рис. 4. Сеть, полученная таким образом из первоначальной сети, будет называться *расширенной сетью*.



Р и с. 4

Легко показать, что необходимые и достаточные условия для существования решения этой задачи с несколькими входами и выходами состоят в следующем:

- 1) сумма V входных потоков равна сумме выходных потоков;
- 2) минимальное сечение в расширенной сети имеет значение V .

Чтобы доказать это, заметим, что необходимость первого условия очевидна, а необходимость второго следует из предположения, что поток в первоначальной сети удовлетворяет условиям задачи. Этот поток можно преобразовать в поток расширенной сети и использовать теорему, из которой вытекает отсутствие сечений со значением, меньшим V . Так как имеются сечения со значением V , проходящие через присоединенные ребра, то значение минимального сечения равно V .

Достаточность этих условий следует из того замечания, что условие 2) благодаря доказанной теореме означает возможность создания в расширенной сети потока слева направо со значением V . Далее по закону Кирхгофа для правого и левого полюсов, используя условие 1), получим, что каждое присоединенное входное и выходное ребро несет поток, равный искомому. Следовательно, это же распределение потока в первоначальной сети решает поставленную задачу.

Ниже дается топологическое решение проблемы раскраски применительно к цветной маркировке проводов в электрических схемах, таких, как релейные панели. В таких схемах имеется некоторое число реле, переключателей и других устройств $A, B, \dots, E$, которые должны быть соединены между собой. Соединительные провода заранее объединены в жгуты таким образом, что все провода, идущие к A , выходят из общего жгута в одной точке, идущие к B , — в другой точке и т. д. Необходимо, чтобы все провода, выводимые из общего жгута в данной точке, имели бы разную окраску, исключаящую возможность ошибочного соединения при монтаже. Произвольную пару устройств может соединить любое число проводов, но один провод может соединять не более двух устройств. Предполагая, что к одному устройству подходит не более m проводов, требуется определить минимальное число цветов, достаточное для выполнения монтажа произвольной схемы.

Т е о р е м а²⁾. Ребра любого графа могут быть окрашены так, что любые два ребра с общим концом будут иметь различные цвета при использовании самое большее $\lceil \frac{3}{2} m \rceil$ цветов, где m — максимальное число ребер, исходящих из одной вершины. Это число цветов необходимо для некоторых графов.

Д о к а з а т е л ь с т в о. Простой граф, требующий $\lceil \frac{3}{2} m \rceil$ цветов, может быть построен следующим образом. Для $m = 2n$ пусть каждая пара из трех вершин A, B, C соединена n ребрами. Так как все ребра, очевидно, должны быть окрашены различно, необхо-

¹⁾ Shannon C., A theorem on colouring the lines of a network, *J. of Math. and Phys.*, 28 (1949), 148.

²⁾ В дальнейшем будет использована терминология, принятая в теории графов [см., например, König D., *Theorie der endlichen und unendlichen Graphen*, Leipzig, 1936, или Кудрявцев Л. Д., О некоторых математических вопросах теории электрических цепей, *Успехи матем. наук*, 3, № 4 (1948), 80]. —Прим. ред.

димо $3n = \left\lfloor \frac{3}{2} m \right\rfloor$ цветов. Если $m = 2n + 1$, то пусть A и B соединены n ребрами, B и C также n ребрами, а A и C соединены $n+1$ ребрами. Здесь все ребра должны быть окрашены различно, и имеется $3n + 1 = \left\lfloor \frac{3}{2} (2n + 1) \right\rfloor = \left\lfloor \frac{3}{2} m \right\rfloor$ ребер. Другой пример для $m = 3$ получается при соединении двух пятиугольников $abcde$ и $ABCDE$ ребрами aA, bD, cB, dE, eC .

Для доказательства достаточности предположим сначала, что m четно. Пусть N — заданный граф; хорошо известно, что можно добавлением ребер и вершин получить из него регулярный граф N' степени m , т. е. такой, в котором каждая вершина является концом ровно m ребер¹⁾. Если можно раскрасить граф N' , то можно, конечно, раскрасить и граф N . Теорема Петерсена²⁾ утверждает, что всякий регулярный граф четной степени $m = 2n$ может быть разложен в n регулярных графов второй степени (факторов)³⁾. Пусть в нашем случае $N_1, N_2, \dots, N_n$ — факторы графа N' ; каждый из них представляет собой совокупность многоугольников (циклов), не касающихся друг друга, и, следовательно, каждый N_i может быть раскрашен самое большее тремя цветами. Это дает в общем $3n = \frac{3}{2} m$ цветов⁴⁾.

Петерсен высказал предположение, что каждый регулярный граф нечетной степени $2n + 1$ без мостиков⁵⁾ может быть разложен в один граф первой и n графов второй степени. Если это предположение верно, то теорема легко доказывается и в случае нечетной степени. Однако это предположение не было доказано для $m > 3$, и здесь воспользуемся другим приемом⁶⁾.

Теорема будет доказываться по индукции; раскрашивание графа N будет производиться в зависимости от раскраски графа, получающегося из N удалением одной вершины с исходящими из нее ребрами. Удалим из N вершину P и $m = 2n + 1$ исходящих из нее ребер (звезду) и предположим, что оставшийся граф может быть раскрашен требуемым способом $3n + 1$ цветами. Пусть вершины, которые были соединены с P в исходном графе, занумерованы числами $1, 2, \dots, s$; и предположим, что было p_1 параллельных ребер

¹⁾ См. цитированную выше книгу Кёнига, стр. 157. — *Прим. ред.*

²⁾ Petersen J., Die Theorie der regulären Graphs, *Acta Mathematica*, 15 (1891), 193. — *Прим. ред.*

³⁾ Доказательство см., например, в цитированной выше книге Кёнига, стр. 161. — *Прим. ред.*

⁴⁾ Это доказательство для четного m было предложено Р. М. Фостером.

⁵⁾ Мостиком называется ребро, не являющееся концевым и не принадлежащее никакому циклу (Кёниг, стр. 179). — *Прим. ред.*

⁶⁾ Конструкция, примененная в доказательстве для случая нечетного m , почти без изменений может быть использована и для четного m .

в первой группе G_1 , соединяющей P с вершиной 1, и т. д. После раскраски оставшегося графа можно использовать для ребер, исходящих из вершины 1, самое меньшее $[(3n + 1) - (2n + 1 - p_1)] = n + p_1$ цветов, для ребер, исходящих из вершины 2, — самое меньшее $n + p_2$ цветов и т. д. Покажем, что при надлежащем выборе этих цветов и подходящей перестановке некоторых цветов в уже раскрашенной части графа всегда можно окрасить ребра, исходящие из P , требуемым образом.

Соберем наши данные в таблицу.

Таблица 1

		цвета									
		1	2	3	.	.	.	.	.	.	$(3n + 1)$
ребра	1	1	1	0	1	.	.	.	.	.	.
	2	1	0	1	0	.	.	.	.	.	.
	.	.	.	.	.	.	.	.	.	.	.
	.	.	.	.	.	.	.	.	.	.	.
	.	.	.	.	.	.	.	.	.	.	.
	$(2n + 1)$	.	.	.	.	.	.	.	.	.	.

В этой таблице $2n + 1$ ребер, исходящих из P , соответствуют строкам, $3n + 1$ цветов — столбцам. Если некоторый цвет имеется для раскраски некоторого ребра, то тогда на пересечении соответствующих столбца и строки ставится 1, в противном случае — 0. В строке, соответствующей ребру из G_i , должно быть $n + p_i$ единиц. Используя описанные ниже три операции, получаем в табл. 1 последовательность единиц вдоль главной диагонали, и это определит способ раскраски графа. Эти операции таковы.

1. Перестановка столбцов. Это соответствует перемене номеров цветов.

2. Перестановка строк. Это соответствует перемене номеров ребер, исходящих из P .

3. Перестановка цветов в цепочке из двух цветов. Две вершины назовем *соединенными цветами* Q и R , если можно перейти от одной из них к другой по цепочке, вдоль которой эти два цвета чередуются. Если имеется правильно (т. е. в соответствии с условиями теоремы) раскрашенный граф и если переставить два цвета в такой цепочке графа вдоль всей ее длины¹⁾ (заметим, что в правильно раскрашенном графе такая цепочка не может разветвляться), то, очевидно, граф останется правильно раскрашенным. Воспользуемся также тем обстоятельством, что если только один из двух определенных цветов встречается около каждой из трех различных вершин, то

¹⁾ Имеется в виду максимальная цепочка, т. е. не являющаяся собственной частью никакой цепочки такого же вида. — *Прим. ред.*

Теперь рассмотрим табл. 1. Предположим, что в первой строке имеется 0 в одном столбце и 1 в другом. Перестановка этих двух цветов в цепочке, исходящей из вершины A , к которой подходит первое ребро, очевидно, эквивалентна перестановке этих двух столбцов в первой строке и во всех остальных строках, которые соответствуют ребрам, подходящим к вершине A . Предположим, что уже получены единицы на главной диагонали D до некоторого места. Покажем, что можно получить следующую единицу на диагонали. Обратимся к табл. 2.

(3n + 1) цветов

$(2n+1)$ ребер

$\gamma \in G_K$

$\beta \in G_J$

$\alpha \in G_I$

$S \quad R \quad T$

Если в строке α имеются единицы, находящиеся в столбце T или правее T , то одна из них может быть перенесена посредством перестановки столбцов на место X . Предположим, что это не имеет места; тогда $n + r_i$ единиц должны находиться левее столбца T в строке α (предполагается, что $\alpha \in G_i$). Следовательно, имеется $n + r_i$ строк над α , имеющих единицы в D в тех же самых столбцах, что и единицы в α . По крайней мере $n + 1$ из этих строк не принадлежат G_i , так как G_i имеет r_i элементов и один из них уже был учтен, а именно α . Пусть β — одна из этих строк, принадлежащая, скажем, G_j . Если β имеет единицу в столбце T или правее T , то перестановкой сначала столбцов, а затем строк α и β можно переместить эту единицу на место X , не затрагивая единиц вдоль D . Предположим, что это не так; тогда $n + r_j$ единиц в строке β рас-

положены левее T и, следовательно, $n + p_j$ строк над строкой α имеют единицы на D в тех же столбцах, что и единицы в β , и из них по крайней мере n не принадлежат G_j (так как G_j имеет только p_j элементов). Далее, над α расположено не более чем $2n$ строк, и поэтому среди n строк, связанных (указанным только что образом) с β , и среди $n + 1$ строк, связанных с α , должна быть по крайней мере одна общая строка, т. е. существует строка, не принадлежащая ни G_i , ни G_j и имеющая единицу в D в том же столбце, в котором имеют единицу α и β . Обозначим эту строку через γ и предположим, что она принадлежит G_k . Если γ имеет единицу в столбце T или правее T , то ее можно переместить на место X путем перестановки сначала столбцов, а затем строк α и γ .

Предположим, что это не имеет места, и нули и единицы на пересечениях α , β , γ и T , S расположены так, как указано на табл. 2, а на месте X стоит 0. Следовательно, по крайней мере одна из строк α , β , γ ¹⁾ не соединена ни с какой из остальных цепочкой с двумя цветами T и S . Если это α , то переменим местами цвета T и S в цепочке, начинающейся в i ; в результате этого единица из пересечения αS перейдет на место X без изменения D . Если это β , то переменим местами цвета в цепочке, начинающейся в j , и переставим строки α и β . Это переместит единицу из пересечения βS на место X , а ее место займет единица из пересечения αS . Если это оказалось γ , то переменим местами цвета в цепочке, начинающейся в k , и переставим строки α и γ так, что единица из γS перейдет на место X , а единица из αS займет ее место.

¹⁾ Точнее, по крайней мере одна из вершин i , j , k .— Прим. ред.

Введение

В известной статье²⁾ А. М. Тьюринг определил класс вычислительных машин, называемых ныне машинами Тьюринга. Можно представить себе, что машина Тьюринга состоит из трех частей: управляющего элемента, считывающей и записывающей головки и бесконечной ленты. Лента разделена на последовательность квадратов, каждый из которых может хранить любой символ из конечного алфавита. Считывающая головка в каждый данный момент воспринимает один квадрат ленты. Она может прочесть записанный там символ и под действием управляющего элемента записать новый символ, а также передвинуться на один квадрат вправо или влево. Управляющий элемент представляет собой устройство с конечным числом внутренних «состояний». В каждый данный момент ближайшая операция машины определяется текущим состоянием управляющего элемента и символом, воспринимаемым считывающей головкой. Эта операция состоит из трех частей: 1) записи нового символа в воспринимаемом квадрате (новый символ, конечно, может совпадать с только что прочитанным); 2) перехода управляющего элемента в новое состояние (которое может совпадать с предыдущим состоянием); 3) движения считывающей головки на один квадрат вправо или влево³⁾.

При подготовке машины к работе на некоторый конечный кусок ленты наносится начальная последовательность символов, а остальная часть ленты оставляется пустой (т. е. заполненной некоторым «пустым» символом). Считывающая головка помещается в некотором начальном квадрате, и машина приступает к вычислениям согласно правилам своей работы. В первоначальной формулировке

¹⁾ Shannon C., A universal Turing mashine with two internal states. Automata Studies, Princeton University Press, 1956, 157.

²⁾ Turing A. M., On computable numbers, with an application to the Entscheidungsproblem, *Proc. London Math. Soc.* (2), 42 (1936), 230—265.

³⁾ Иногда предполагают, что считывающая головка может также оставаться на месте.— *Прим. перев.*

Тьюринга квадраты через один предназначались для записи окончательного ответа и для промежуточных вычислений. Эта и другие подробности первоначального определения были изменены в позднейших изложениях теории.

Тьюринг показал, что возможно построить универсальную машину, способную работать как любая частная машина Тьюринга, если ее снабдить описанием этой частной машины. Описание наносится на ленту универсальной машины по определенному коду, подобно начальной последовательности кодов частной машины. Тогда универсальная машина воспроизводит работу частной.

Наша главная цель — показать, что можно построить универсальную машину Тьюринга, использующую одну ленту и имеющую лишь два внутренних состояния. Будет показано также, что с одним внутренним состоянием этого сделать нельзя. В заключение этой статьи дается построение универсальной машины Тьюринга только с двумя символами на ленте.

Универсальная машина Тьюринга с двумя состояниями

В общих чертах метод построения машины таков. Для произвольной машины Тьюринга A с алфавитом из m букв (символов, записываемых на ленте, включая пустой квадрат) и с n внутренними состояниями строится машина B с двумя внутренними состояниями и алфавитом не более чем из $4mn + m$ символов. Машина B будет работать по существу так же, как и машина A . Во всех квадратах, кроме символа, воспринимаемого считывающей головкой и смежного с ним, на ленте машины B записано то же, что и на ленте машины A в соответствующие моменты работы двух машин. Если в качестве A выбрать универсальную машину Тьюринга, то и B будет универсальной машиной Тьюринга.

Машина B моделирует поведение машины A , но хранит информацию о внутреннем состоянии машины A с помощью символов, записанных в квадрате под считывающей головкой и в квадрате, к которому считывающая головка машины A собирается перейти. Основная задача — своевременно освежать эту информацию и держать ее под считывающей головкой. Если последняя передвигается, то информацию о состоянии надо перенести в новый квадрат, используя всего два внутренних состояния машины B . Пусть, например, следующим состоянием машины A должно быть состояние 17 (согласно какой-нибудь системе счисления). Чтобы перенести его символ, считывающая головка «качается» вперед — назад между старым и новым квадратом 17 раз (точнее 18 раз в новый квадрат и 17 назад, в старый). В течение этого процесса символ, стоящий в новом квадрате, пробегает своего рода последовательность счета, которая оканчивается символом, соответствующим

состоянию 17 и в то же время сохраняющим информацию о предыдущем символе в этом квадрате. Процесс качания возвращает также старый квадрат к одному из элементарных символов (находящихся во взаимнооднозначном соответствии с символами, используемыми машиной A), а именно к тому элементарному символу, который должен быть записан в этом квадрате после окончания этой операции.

Формально машина B строится так. Пусть символы алфавита машины A суть $A_1, A_2, \dots, A_m$ и пусть ее состояния суть $S_1, S_2, \dots, S_n$. В машине B поставим в соответствие алфавиту машины A m элементарных символов $B_1, B_2, \dots, B_m$. Затем определим $4mn$ новых символов, соответствующих парам из состояния и символа машины A , снабженных двумя новыми двузначными индексами. Эти символы обозначим через $B_{i,j,x,y}$, где $i = 1, 2, \dots, m$ (соответственно символам), $j = 1, 2, \dots, n$ (соответственно состояниям), $x = +$ или $-$ (в зависимости от того, передает или получает информация квадрат ленты во время качания) и $y = R$ или L (в зависимости от того, вправо или влево от воспринимаемого квадрата должна передвинуться считывающая головка при качании).

Два состояния машины B назовем α и β . Эти состояния используются двояко. Во-первых, при первом шаге качания они переносят в ближайший очередной квадрат информацию о том, вправо (α) или влево (β) от нового квадрата лежит старый. Эта информация нужна в новом квадрате, чтобы управляющий элемент передвинул считывающую головку назад в нужном направлении. После первого шага информация об этом сохраняется в новом квадрате с помощью записанного там символа (последний индекс y). Во-вторых, состояния α и β используются, чтобы сообщить из старого квадрата в новый о конце качания. После первого шага качания состояние β переносится в новый квадрат вплоть до конца качания, когда переносится α . Это означает конец операции, и новый квадрат начинает затем действовать как передатчик и управляет следующим шагом вычисления.

Описывая машину B , указывают, что именно она выполняет при чтении произвольного символа в произвольном состоянии. Выполняет же она три операции: записывает новый символ, переходит в новое состояние и передвигает считывающую головку вправо или влево. Таблица работы машины B такова:

Символ	Состояние $\rightarrow$ символ	Состояние	Направление	
B_i	$\alpha \rightarrow B_{i,1,-,R}$	α	R	$(i=1, 2, \dots, m)$ (1)
B_i	$\beta \rightarrow B_{i,1,-,L}$	α	L	$(i=1, 2, \dots, m)$ (2)

Продолжение

Символ	Состояние $\rightarrow$ символ	Состоя- ние	Нап- равле- ние	
$B_{i,j,-,x}$	$\beta \rightarrow B_{i;(j+1),-,x}$	α	x	$(i=1, 2, \dots, m)$ $(j=1, 2, \dots, n-1)$ $(x=R, L)$
$B_{i,j,+,x}$	α или $\beta \rightarrow B_{i;(j-1),+,x}$	β	x	$(i=1, 2, \dots, m)$ $(j=2, \dots, n)$ $(x=R, L)$
$B_{i,1,+,x}$	α или $\beta \rightarrow B_i$	α	x	$(i=1, 2, \dots, m)$ $(x=R, L)$

Эти операции пока что никак не зависят от таблицы работы машины A (кроме числа используемых символов). Операции же следующего и последнего типа формулируются в терминах таблицы работы моделируемой машины. Предположим, что машина A имеет формулу операций

$$A_i; S_j \rightarrow A_k; S_i; \begin{matrix} R \\ L \end{matrix}. \quad (6)$$

Тогда по определению машина B имеет формулу

$$B_{i,j,-,x}; \alpha \rightarrow B_{k,l,+, \begin{matrix} \beta \\ \alpha \end{matrix}}; \begin{matrix} R \\ L \end{matrix}, \quad (7)$$

причем, если в формуле (6) употреблена верхняя буква (R), то верхние буквы употребляются также и в формуле (7) и обратно.

Проследим цикл работы системы, состоящий из одной операции машины A и соответствующей серии операций машины B .

Пусть машина A читает символ A_3 и находится в состоянии S_7 и пусть ее таблицей работы предусмотрена запись символа A_8 , переход в состояние S_4 и движение вправо. Машина B читает (по индуктивному предположению) символ $B_{3,7,-,x}$ [значение x (R или L) зависит от предыдущих операций и безразлично для дальнейшего]. Машина B будет находиться в состоянии α . В соответствии с формулой (7) машина B запишет $B_{8,4,+,R}$, перейдет в состояние β и передвинется вправо. Предположим, что квадрат справа в машине A содержит A_{13} ; соответствующий квадрат в машине B содержит B_{13} . Приходя в этот квадрат в состоянии β , машина B в силу формулы (2) записывает $B_{13,1,-,L}$, переходит в состояние α и движется назад, влево. Это служит началом переноса информации о состоянии с помощью процесса качания. Приходя в квадрат слева, машина B читает $B_{8,4,+,R}$ и в силу формулы (4) записывает $B_{8,3,+,R}$, переходит в состояние β и передвигается опять

вправо. Затем в силу формулы (3) она записывает $B_{13, 2, -, L}$, переходит в состояние α и возвращается налево. Весь ход процесса приведен в табл. 1.

Таблица 1

Символ в левом квадрате	Состояние	Символ в правом квадрате
$B_{3, 7, -, x}$	β	B_{13}
$B_{8, 4, +, R}$	α	
	β	$B_{13, 1, -, L}$
$B_{8, 3, +, R}$	α	
	β	$B_{13, 2, -, L}$
$B_{8, 2, +, R}$	α	
	β	$B_{13, 3, -, L}$
$B_{8, 1, +, R}$	α	
	α	$B_{13, 4, -, L}$

Указанные операции завершают перенос информации о состоянии в квадрат справа и выполнение команды, отданной квадратом слева. В квадрате слева записан теперь символ B_8 (соответствующий символу A_8 в машине A), а в квадрате справа — символ $B_{13, 4, -, L}$; считывающая головка находится на этом квадрате, имея внутреннее состояние α . Это возвращает нас к ситуации, аналогичной той, которая имела место в начале рассмотренного цикла. Рассуждая по индукции, видим, что машина B моделирует поведение машины A .

Чтобы заставить машину B работать аналогично машине A , заполним начальную ленту машины B соответственно начальной ленте машины A (с заменой A_i на B_i), за исключением квадрата, занимаемого считывающей головкой в начальный момент. Если S_j — начальное состояние машины A и A_i — начальный символ в этом квадрате, то в соответствующем квадрате ленты машины B записываем $B_i, j, -, R$ (или L) и приводим B в состояние α .

Невозможность создания универсальной машины Тьюринга с одним состоянием

Теперь покажем, что нельзя построить универсальную машину Тьюринга, использующую одну ленту и только одно внутреннее состояние.

Допустим, что такая машина существует. Если записать подходящее «описательное число» конечной длины на куске ленты (оставив прочую часть ленты пустой) и поместить считывающую

головку в соответствующий квадрат, то машина будет вычислять любое вычислимое число и, в частности, вычисляемые иррациональные числа, например $\sqrt{2}$. Покажем, что это невозможно.

Согласно первоначальному замыслу Тьюринга, при вычислении $\sqrt{2}$ последовательные знаки $\sqrt{2}$ (скажем, в двоичном представлении) записываются машиной на определенной последовательности квадратов ленты (например, на четных квадратах, тогда как нечетные предназначаются для промежуточных вычислений). Приводимое ниже доказательство исходит из вычисления $\sqrt{2}$ в таком виде, хотя будет ясно, что, изменив доказательство, можно принять во внимание и другие разумные интерпретации «вычисления $\sqrt{2}$ ».

Ввиду иррациональности $\sqrt{2}$ двоичные знаки его разложения не станут повторяться периодически. Поэтому, если показать для машины с одним состоянием, что или 1) во всех ее квадратах, кроме конечного числа, будет в конце концов записан один и тот же символ, или 2) во всех квадратах, кроме конечного числа, символы будут без конца сменяться, то тем самым будет получен желаемый результат.

Сначала возьмем ленту, пустую и бесконечную в обе стороны от описательного числа для $\sqrt{2}$. Если считывающая головка приходит на пустой квадрат, то она должна или остановиться здесь на все время, или же передвинуться вправо либо влево. Так как имеется лишь одно состояние, то ее поведение не зависит от предшествовавшего вычисления. В первом случае считывающая головка никогда не отойдет дальше, чем на один квадрат, от описательного числа, и вся лента, кроме конечного отрезка, будет постоянно пуста. Если же головка передвигается влево с некоторого пустого квадрата, то или левая бесконечная часть пустой ленты не участвует в вычислении и поэтому не требует рассмотрения, или, если она участвует в вычислении, считывающая головка с этого времени непрерывно движется влево, записывая во всех квадратах, ранее пустых, один и тот же символ. Следовательно, лента становится одинаковой всюду слева от некоторого конечного отрезка и пустой справа от него и не может хранить записи числа $\sqrt{2}$. Аналогичное положение возникает при движении головки вправо от первоначально пустого квадрата. Поэтому двусторонне бесконечная лента ничуть не лучше односторонней, и можно из соображений симметрии предположить, что лента односторонне бесконечна вправо от описательного числа.

Теперь рассмотрим следующую операцию. Поместим считывающую головку в первом квадрате бесконечной пустой полосы¹⁾. Машина будет вычислять некоторое время, и, быть может, считы-

¹⁾ Справа от описательного числа. — *Прим. перев*

вающая головка переместится назад из этой полосы по направлению к описательному числу. Тогда возвратим ее снова в первый квадрат первоначально пустой полосы. Если головка опять переместится к описательному числу, то опять поместим ее в этом первом квадрате и т. д. Число раз, которое головку можно поместить таким образом над этим первым квадратом, будет называться числом отражений машины и обозначаться через R . Оно равно или целому числу $1, 2, 3, \dots$ или ∞ .

Теперь поместим считывающую головку в квадрат описательного числа, являющийся начальным для предполагаемого вычисления $\sqrt{2}$. Спустя некоторое время считывающая головка, возможно, выйдет из части ленты, где записано описательное число. Возвратим ее в последний квадрат описательного числа. Спустя некоторое время она, возможно, опять выйдет. Продолжим этот процесс до тех пор, пока возможно. Число выходов головки равно или целому числу $0, 1, 2, 3, \dots$ или ∞ . Это число S назовем числом отражений для данного описания $\sqrt{2}$.

Если S конечно и R (может быть ∞) $\geq S$, то спустя конечное время считывающая головка застрянет в части ленты, где первоначально находилось описательное число. Изменится лишь конечная часть пустой ленты, и машина не вычислит $\sqrt{2}$.

Если как R , так и S бесконечны, то считывающая головка будет возвращаться неограниченное число раз в часть ленты с описательным числом. Экскурсий в первоначально пустую часть будут или ограничены по длине, или нет. Если они ограничены, то изменится лишь конечная часть пустой ленты, как и в предыдущем случае. Если же нет, то вся лента, кроме конечной части, будет обрабатываться считывающей головкой неограниченное число раз. Так как имеется только одно состояние и алфавит символов конечен, то символ, записываемый в квадрате, который посещается неограниченное число раз, должен или стать постоянным (одним и тем же для всех таких квадратов), или изменяться циклически бесконечное число раз. В первом случае вся первоначально пустая лента становится одинаковой и не может изображать $\sqrt{2}$. Во втором случае она непрерывно изменяется и не может изображать никакого результата вычисления¹⁾.

¹⁾ Рассуждения автора здесь не вполне ясны, и, кроме того, он использует условие конечности алфавита символов не по существу. Можно дать другое доказательство без предположения, что число символов конечно.

Пусть вся лента, кроме конечной ее части, обрабатывается головкой неограниченное число раз. А priori возможны два случая: либо во всяком квадрате, кроме конечного их числа, символы сменяются неограниченное число раз (и тогда, очевидно, запись на ленте не может представлять результат вычисления $\sqrt{2}$), либо это не так. Во втором случае сколь угодно далеко

Если $R < S$, то считывающая головка в конце концов уйдет в первоначально пустую часть ленты и там останется. Можно показать, что в этом случае символы в первоначально пустой части становятся, за исключением конечного их числа, одинаковыми. В самом деле, либо головка передвигается вправо из первого пустого квадрата во второй пустой квадрат по меньшей мере $R + 1$ раз, либо нет. В последнем случае спустя конечное время считывающая головка остается в бывшем первом пустом квадрате, и вся лента, кроме конечной части, остается пустой. Если же головка уходит вправо $R + 1$ раз, то она не вернется в первый квадрат, первоначально пустой, ибо R есть число отражений для пустой ленты. В этом первом квадрате будет записан результат $2R + 1$ действий над пустым квадратом (при $R + 1$ приходах головки слева и R приходах справа). Во втором первоначально пустом квадрате в конце концов будет записан тот же самый постоянный символ, потому что ко второму квадрату применимо такое же рассуждение, как и к первому. Во всех случаях машина работает на той же самой ленте (бесконечном ряде пустых квадратов) и приходит одно и то же число раз справа (R) и соответственно ($R + 1$) слева. Тем самым все возможные случаи исчерпаны и доказательство завершено.

Моделирование машины Тьюринга с помощью только двух символов на ленте

Теперь покажем, что можно также построить машину C , работающую подобно любой заданной машине Тьюринга A и использующую только два символа 0 и 1 на своей ленте; один из них, например 0, служит символом пустого квадрата. Пусть по-прежнему заданная машина A имеет m символов и n внутренних состоя-

найдется квадрат q , обладающий тем свойством, что, начиная с некоторого момента, символ в q не изменяется (обозначим этот символ через C); тогда q можно выбрать так, чтобы слева от q имелся некоторый квадрат q^* , в который головка возвращается неограниченное число раз.

Пусть головке предписано сдвинуться вправо по прочтении символа C . Тогда головка, попав на квадрат q , после того как в нем постоянно записывается символ C , не сможет, начиная с этого момента, попасть в q^* . Аналогично, если по прочтении символа C головка должна сдвинуться влево, то она не сможет, начиная с некоторого момента, попасть в квадраты, лежащие правее q .

Это говорит о том, что второй из двух а priori возможных случаев не может иметь места. Итак, возможен лишь первый случай, при котором машина не вычислит $\sqrt{2}$.

Так как в остальных частях доказательства Шеннон нигде не опирается на условие конечности алфавита символов, то можно считать его теорему справедливой также для «машин», работающих с бесконечным числом символов (и одним состоянием). — *Прим. перев.*

ний. Пусть l — наименьшее целое число, для которого $m \leq 2^l$. Тогда символам машины A можно сопоставить двоичные последовательности длины l таким образом, что различным символам будут соответствовать различные последовательности. При этом пустому символу машины A ставится в соответствие последовательность из l нулей. Машина C будет работать с двоичными последовательностями. Элементарная операция машины A будет соответствовать в машине C переходу считывающей головки на $l - 1$ квадратов вправо (с сохранением считанной информации во внутреннем состоянии головки), затем обратному переходу на $l - 1$ квадратов влево, записи новых символов по пути и, наконец, движению вправо или влево на l квадратов в соответствии с движением считывающей головки машины A . В течение этого процесса состояние машины A , конечно, сохраняется и в машине C . Замена старого состояния новым происходит в конце операции считывания.

Формально машина C строится так. Состояниям $S_1, S_2, \dots, S_n$ машины A ставятся в соответствие состояния $T_1, T_2, \dots, T_n$ машины C (последние будут встречаться, когда машина C начинает операцию, считывая первый символ в двоичной последовательности длины l). Для каждого из этих T_i определим два состояния T_{i0} и T_{i1} . Если машина C находится в состоянии T_i и читает символ 0, то она движется вправо и переходит в состояние T_{i0} . Если она читает 1, то движется вправо и переходит в состояние T_{i1} . Таким образом, с помощью этих двух состояний машина запоминает, каким был первый символ двоичной последовательности. Для каждого из этих T_{i0}, T_{i1} определим опять по два состояния: T_{i00}, T_{i01} и T_{i10}, T_{i11} . Если, например, машина находится в состоянии T_{i0} и читает символ 0, то она переходит в состояние T_{i00} и т. д. Таким образом, с помощью этих состояний запоминается начальное состояние и два первых символа, прочитанных в процессе работы машины. Продолжим такое построение состояний вплоть до $l - 1$ шагов, получив в итоге $(2^l - 1)n$ состояний. Эти состояния можно обозначить через

$$T_{i, x_1, x_2, \dots, x_s}; \quad i = 1, 2, \dots, n; \quad x_j = 0, 1; \quad s = 0, 1, \dots, l - 1.$$

Если машина находится в одном из этих состояний ($s < l - 1$) и читает 0 или 1, то она движется вправо и 0 или 1 делается следующим индексом состояния. Когда $s = l - 1$, машина читает последний символ последовательности длины l . Теперь правила операций зависят от конкретных правил машины A . Определим два новых множества состояний, которые несколько похожи на введенное выше множество состояний T , но соответствуют не считыванию, а записи:

$$R_{i, x_1, x_2, \dots, x_s} \quad \text{и} \quad L_{i, x_1, x_2, \dots, x_s}.$$

Последовательность $x_1, x_2, \dots, x_{l-1}, x_l$ соответствует некоторому символу машины A . Предположим, что, если машина A читает этот символ и находится в состоянии i , то она записывает символ, соответствующий двоичной последовательности $y_1, y_2, \dots, y_{l-1}, y_l$, переходит в состояние j и движется, скажем, вправо. Тогда по определению машина C , будучи в состоянии $T_{i, x_1, x_2, \dots, x_{l-1}}$ и читая символ x_i , переходит в состояние $R_{j, y_1, y_2, \dots, y_{l-1}}$, записывает символ y_l и движется влево. В любом из состояний $R_{i, y_1, y_2, \dots, y_s}$ (или $L_{i, y_1, y_2, \dots, y_s}$) машина C записывает y_s , движется влево и переходит в состояние $R_{i, y_1, y_2, \dots, y_{s-1}}$ (или $L_{i, y_1, y_2, \dots, y_{s-1}}$). Посредством этого процесса двоичная последовательность, соответствующая новому символу, записывается вместо старой двоичной последовательности. При $s = 1$ эта запись заканчивается на символе y_1 . Остается только передвинуть считывающую голову на l квадратов вправо или влево в зависимости от того, находится ли машина в состоянии R или в состоянии L . Это делается с помощью множеств состояний U_{is} и V_{is} ($i = 1, 2, \dots, n$; $s = 1, 2, \dots, l - 1$). В состоянии R_{ix_1} машина записывает x_1 , движется вправо и переходит в состояние U_{i1} . В каждом из состояний U машина продолжает движение вправо, не записывая ничего и переходя в состояние U со следующим по величине индексом, пока не будет достигнуто последнее состояние U . Таким образом, U_{is} вызывает движение вправо и состояние $U_{i, s+1}$ ($s < l - 1$). Наконец, состояние $U_{i, l+1}$ приводит—после движения вправо—к состоянию T_i , завершая тем самым цикл. Аналогично, L_{ix_1} приводит к движению влево и состоянию V_{i1} ; V_{is} дает движение влево и $V_{i, s+1}$ ($s < l - 1$); наконец $V_{i, l-1}$ дает движение влево и T_i .

Начальная лента машины C представляет собой, конечно, начальную ленту машины A , где каждый символ замещен соответствующей ему двоичной последовательностью. Если работа машины A начинается с какого-то определенного символа, то работа машины C начнется с самого левого двоичного символа соответствующей группы; если машина A начинает работу в состоянии S_i , то C начнет работу в состоянии T_i .

Машина C имеет самое большее $n(1 + 2 + 4 + \dots + 2^{l-1}) = n(2^l - 1)$ состояний T , самое большее $n(2^l - 2)$ состояний R и $n(2^l - 2)$ состояний L и, наконец, $2n(l - 1)$ состояний U и V . Таким образом, всего требуется не более чем $3n2^l + n(2l - 7)$ состояний. Так как $2^l < 2m$, то эта верхняя оценка числа состояний меньше, чем $6mn + n(2l - 7)$, что в свою очередь заведомо меньше, чем $8mn$.

Полученные нами результаты вместе с интуитивными соображениями подсказывают, что можно заменять символы состояниями и обратно (в некоторых пределах), не изменяя намного произведения mn . При переходе к двум состояниям произведение возрастает

приблизительно в восемь раз. При переходе к двум символам произведение возрастает приблизительно в шесть раз, но не более чем в восемь раз. Эти «коэффициенты потерь» — 6 и 8 — вызваны, вероятно, нашим методом микромоделирования, при котором каждая элементарная операция машины A моделируется в машине B . Если конструировать машину B , заботясь лишь о том, чтобы она имела те же вычислительные возможности в широком смысле, что и машина A , то произведение mn изменилось бы значительно меньше. Во всяком случае, число логических элементов (например, реле), необходимых для физической реализации машины, является небольшим постоянным кратным (приблизительно двукратным в случае реле) двоичного логарифма произведения mn , и потому коэффициент 6 или 8 потребует лишь немногих дополнительных реле.

Нахождение минимального произведения mn для универсальной машины Тьюринга¹⁾ представляет собой интересную нерешенную задачу.

¹⁾ В настоящее время наиболее сильный результат в этом направлении получил С. Ватанабе. Он построил универсальную машину Тьюринга с $m = 5$ и $n = 8$ (Watanabe S., 5-symbol 8-state and 5-symbol 6-state universal Turing machines, *J. Assoc. Comput. Mach.*, 8, № 4 (1941), 476; Русский перевод см. Кибернетический сборник, вып. 6, ИЛ, 1963, 80). — *Прим. ред.*

ВЫЧИСЛИМОСТЬ НА ВЕРОЯТНОСТНЫХ МАШИНАХ¹⁾

В статье рассматривается следующий вопрос: имеются ли задачи, которые может решать машина со случайным элементом, но не может решать детерминистская машина?

Конечно, в таком виде вопрос выглядит слишком неопределенным, чтобы быть доступным для математического анализа. В дальнейшем ограничим его в двояком отношении. Будет дано точное определение класса рассматриваемых машин и столь же точное определение выполняемых ими задач. Ясно, что наши выводы будут сильно зависеть от выбора этих двух определений, и поэтому их нельзя считать полным решением первоначального неформального вопроса. Читатель должен иметь в виду, что полученные нами результаты неприменимы вне области, обусловленной выбранными определениями. В частности, они указывают на возможность перечисления *бесконечных* множеств и вычисления *бесконечных* последовательностей, но не дают информации о типе вероятностных машин, предназначенных для решения конечных задач, например об относительной сложности и относительном быстродействии таких машин.

Подобные затруднения присутствуют неявно при всяком применении математики к некоторому поставленному неформально вопросу. Формулировка такого вопроса в точной математической форме неизбежно сосредоточивает внимание на некоторых его сторонах, в то время как другие стороны, одинаково и даже более важные в иной ситуации, могут оказаться совершенно неучтенными.

Основной текст настоящей статьи содержит определения, примеры и формулировки результатов. Доказательства перенесены в приложение, так как большей частью они представляют собой более или менее сложные построения, которые не являются абсолютно необходимыми для понимания работы.

¹⁾ De Leeuw K., Moore E. F., Shannon C. and Shapiro N., Computability by probabilistic machines, Automata studies, Princeton University Press, 1956, 183—212.

Читатели, знакомые с теорией рекурсивных функций, без труда увидят, что наши теоремы принадлежат этой теории и легко могут быть переведены на ее язык. Ввиду этого доказательства получают двойное значение. Если считать определенным интуитивное понятие *эффективного процесса*¹⁾, то их можно рассматривать как полные доказательства. С другой стороны, их можно рассматривать как *указания*, исходя из которых, можно теоремы, сформулированные на языке теории рекурсивных функций, доказать формально средствами этой теории. Такие формальные доказательства не приводятся, чтобы не отвлекать читателя от содержания работы. Однако возможность формализации ясна всякому, кто знаком с теорией рекурсивных функций.

Вероятностные машины в сравнении с детерминистскими

Начнем с точного определения некоторого класса вычислительных машин. Эти машины имеют вход и выход. Входом является бесконечная последовательность двоичных знаков, которая может подаваться двумя различными способами: 1) как фиксированная последовательность (физически можно представить заранее подготовленную бесконечную ленту); 2) как выход двоичного случайного устройства (с вероятностью p появления 1). В последнем случае мы имеем вероятностную машину. Затем очертим точный класс задач, подлежащих решению на этих машинах, а именно перечисление с положительной вероятностью множеств выходных символов. После этого ключевой результат настоящей статьи — теорема 1 — используется для получения подхода к ответу на наш вопрос. Ответ дается теоремой 2. Он гласит, что если случайное устройство имеет вероятность p , где p — *вычислимое действительное число* (т. е. действительное число, для которого существует эффективный процесс нахождения любого знака в его двоичном разложении), то любое множество, перечислимое посредством вероятностной машины рассматриваемого типа, перечислимо посредством детерминистской машины. Это не справедливо, если p — невычислимое действительное число. Аналогичные результаты получаются для случая, когда учитывается порядок следования выходных символов. Положение подытоживается теоремой 3.

Представим наши машины в виде объектов, которые воспринимают в качестве входа ленту с записанными на ней нулями и единицами и выдают в качестве выхода ленту, на которой они могут записывать конечную или счетную последовательность выходных

¹⁾ Об эффективном процессе см. в [9] или [2]. Советуем читателю, не знакомому с понятием эффективного процесса, прочесть одну из этих работ, прежде чем идти дальше. Эффективные процессы рассматриваются также в [1], [7] и [8].

символов $s_1, s_2, s_3, \dots$ (этими символами могут быть конфигурации, образованные из некоторого конечного алфавита). Предполагается, что машина имеет некоторое начальное состояние, в которое она приводится перед началом любой операции (подобно тому, как счетно-клавишная машина гасится перед началом вычислений). Другое разумное требование к машине таково: должна существовать эффективная методика, позволяющая определить, какова будет последовательность выходных символов $(s_{j_1}, \dots, s_{j_r})$, если в начальном состоянии машине предъявлена входная последовательность нулей и единиц $(a_1, \dots, a_n)$. Эту выходную последовательность обозначим через $f(a_1, \dots, a_n)$. Так как нас интересует только зависимость между входом и выходом машин, то можно отвлечься от их внутреннего устройства и принять это требование за определение. (Хотя этим дано абстрактное определение, дальше будем говорить о конкретных машинах.)

О п р е д е л е н и е. *Машина* есть функция, которая для каждого n ставит в соответствие каждой конечной последовательности из n нулей и единиц $(a_1, \dots, a_n)$ некоторую конечную последовательность $f(a_1, \dots, a_n) = (s_{j_1}, \dots, s_{j_r})$, состоящую из элементов заданного множества S , причем выполняются условия:

1) $f(a_1, \dots, a_n)$ есть начальный отрезок в

$$f(a_1, \dots, a_n, \dots, a_{n+m}),$$

если $(a_1, \dots, a_n)$ есть начальный отрезок в

$$(a_1, \dots, a_n, \dots, a_{n+m});$$

2) f — вычислимая функция, т. е. существует эффективный процесс¹⁾, позволяющий определить $f(a_1, \dots, a_n)$ по заданной $(a_1, \dots, a_n)$.

Это определение распространяется на бесконечные последовательности следующим образом: если $A = (a_1, a_2, a_3, \dots)$, то $f(A)$ есть последовательность с начальными отрезками $f(a_1, \dots, a_n)$.

Неформально работу машины можно представить себе следующим образом. Получив на входе $(a_1, \dots, a_n)$, она воспринимает a_1 и записывает последовательность символов $f(a_1)$ на выходе; затем она воспринимает a_2 и записывает символы на ленте после $f(a_1)$, производя $f(a_1, a_2) \dots$; затем она воспринимает a_n и записывает символы на ленте после $f(a_1, \dots, a_{n-1})$, производя

$$f(a_1, \dots, a_{n-1}, a_n)$$

и, наконец, останавливается.

Теперь дадим несколько конкретных примеров машин. Этим преследуется двоякая цель: проиллюстрировать наше понятие машины и подготовить материал для иллюстрации новых

¹⁾ См. предыдущее примечание.

понятий, вводимых далее. Пока можно при чтении пропустить ряд примеров.

Машина № 1. Выходными символами этой машины являются упорядоченные пары целых чисел (a, r) . Для каждого входного символа a машина записывает выходной символ (a, r) , если a был r -м входным символом на входной ленте. Иными словами, $f(a_1, \dots, a_n) = ((a_1, 1), (a_2, 2), \dots, (a_n, n))$.

Машина № 2. Пусть g — целочисленная положительная функция и пусть она вычислима, т. е. существует эффективный процесс для нахождения $g(n)$ по заданному n . Пусть машина записывает символ $[r, g(r)]$ после восприятия r -го входного символа. В этом случае $f(a_1, \dots, a_n) = [(1, g(1)), (2, g(2)), \dots, (n, g(n))]$ и выход не зависит от входа. Эта машина и машина № 1 представляют крайние случаи. Машина № 2 забывает о входе, а № 1 по существу копирует вход на выходной ленте. Заметим, что объект № 2 не был бы машиной в смысле нашего определения, если бы функция g не была вычислима, ибо в этом случае не существовало бы эффективного процесса, позволяющего определить выход по заданному входу.

Машина № 3. Пусть машина записывает символ 0, встретив первый нуль во входной последовательности. В остальное время она не записывает ничего. Тогда $f(a_1, \dots, a_n) = (0)$, если один из символов a_i является нулем. В противном случае машина не записывает ничего, и такая последовательность будет «пустой последовательностью» $(\cdot)$; $f(1, \dots, 1) = (\cdot)$.

Машина № 4. Пусть $f(a_1, \dots, a_n) = (a_1)$. Машина просто копирует на выходной ленте первый входной символ, не записывая больше ничего.

Машина № 5. Пусть машина записывает выходной символ r , если максимальная длина групп единиц, воспринятых машиной, равна r . Например, $f(1) = (1)$, $f(1, 0) = (1, 1)$, $f(1, 0, 1) = (1, 1, 1)$, $f(1, 0, 1, 1) = (1, 1, 1, 2)$.

Машина № 6. Пусть $h(r, s)$ — вычисляемая целочисленная положительная функция. Машина не записывает ничего, пока не встретит первой единицы на входе; если эта единица — первый входной символ, то машина вообще ничего не печатает; если же эта единица встретилась после r нулей, то записывается символ $[1, h(r, 1)]$. После следующего входного символа машина записывает символ $[2, h(r, 2)]$, после следующего — символ $[3, h(r, 3)]$ и т. д. Пусть функция $h_r(s)$ с фиксированным r определяется соотношением $h_r(s) = h(r, s)$. Тогда работа машины фактически состоит в вычислении функции h_r при поступлении в машину начальной группы нулей и следующей за ними единицы.

Машина № 7. Эта машина задается соотношением $f(a_1, \dots, a_n) = [(1, r_1), (2, r_2), \dots, (q, r_q)]$, где целое число r_i

получается следующим образом. Возьмем первые 100_s знаков в $(a_1, \dots, a_n)$. Пусть p_s — доля знаков, равных 1. Тогда r_s есть s -й знак в двоичном разложении числа p_s . Число q есть наибольшее s , такое, что $100^s \leq n$. Аналогичная машина используется при доказательстве части теоремы 1.

Хотя работа машины определена только для конечных входных последовательностей, тем не менее ясно, что сделала бы такая машина, получая бесконечную входную последовательность $A = (a_1, a_2, a_3, \dots)$ и обладая бесконечной выходной лентой для записи. Работа машины определена тем фактом, что известен ее выход для всех конечных входных последовательностей. Таким образом, машина ставит в соответствие каждой бесконечной входной последовательности выходную последовательность, которая может быть или не быть бесконечной.

Если вводить в машину № 1 бесконечную ленту, на которой записаны только единицы, то выходная последовательность будет $[(1, 1), (1, 2), (1, 3), \dots]$. Для этой же ленты машина № 3 не даст никакого выхода [пустая последовательность $(\cdot)$], а машина № 4 даст выходную последовательность (1).

Рассмотрим множество символов, появляющихся на выходной ленте машины, когда машина получает некоторую бесконечную входную последовательность A . (Если выходная последовательность есть $(1, 1, 1, \dots)$, то это множество состоит лишь из одного символа 1, а если наша выходная последовательность есть $(1, 3, 5, 7, \dots)$, то множество символов состоит из всех нечетных чисел.) Таким образом, машина ставит в соответствие каждой входной последовательности A некоторое множество выходных символов, которое она перечисляет. (Не будем здесь обращать внимание на порядок, в котором это множество перечисляется, а также на то, перечисляется ли оно с повторениями или без повторений. Это будет сделано позже, при рассмотрении вычисления последовательностей.)

Например, машина № 1 ставит в соответствие любой бесконечной входной последовательности $A = (a_1, a_2, \dots)$ множество выходных символов, состоящее из всех (a_n, n) , где n пробегает целые положительные числа. Машина № 2 ставит в соответствие любой входной последовательности множество всех $[n, g(n)]$, где n пробегает целые положительные числа. Машина № 3 ставит в соответствие входной последовательности, состоящей только из единиц, пустое множество. Такому же входу машина № 5 ставит в соответствие множество всех целых положительных чисел.

О п р е д е л е н и е. Машина, получающая бесконечную входную последовательность $A = (a_1, a_2, a_3, \dots)$, будет называться A -машиной. Будем говорить, что она A -перечисляет множество выходных символов, записываемых ею. Множество каких-либо сим-

волов называется *A-перечислимым*, если существует *A-машина*, которая *A-перечисляет* это множество.

A-перечислимые множества, где *A* есть последовательность, состоящая только из единиц, обычно называются *рекурсивно-перечислимыми* множествами. Мы назовем их *1-перечислимыми* множествами. *A-машина* будет называться *1-машиной*, если *A* состоит целиком из единиц.

Каждой бесконечной последовательности $A = (a_1, a_2, \dots)$ можно поставить в соответствие множество всех пар (n, a_n) , где a_n есть n -й элемент в последовательности *A*, а n пробегает все целые положительные числа. Это множество обозначим через A' . Последовательность *A* называется *вычислимой*, если существует эффективный процесс, с помощью которого можно определить n -й член в *A* для каждого n . Три следующие леммы доказываются в приложении, хотя их доказательства почти очевидны.

Л е м м а 1. *Последовательность A вычислима тогда и только тогда, когда соответствующее множество A' является 1-перечислимым.*

Л е м м а 2. *Если A — вычислимая последовательность, то любое A-перечислимое множество 1-перечисливо (и обратно).*

Л е м м а 3. *Если A — невычислимая последовательность, то существуют множества A-перечислимые, но не 1-перечислимые.*

Лемма 3 и существование невычислимых последовательностей будут важны для нас позже.

Теперь желательно присоединить к машине случайное устройство, чтобы построить вероятностную машину. Если у нас есть устройство, записывающее нули и единицы на ленте, и если единицы появляются с вероятностью p , нули — с вероятностью $1 - p$ ($0 < p < 1$), причем каждая запись независима от предыдущих, то выход такого устройства можно использовать как входную ленту машины. Такая комбинация случайного устройства и машины будет называться *p-машиной*.

Поставленный во введении вопрос о существовании вероятностных машин, более общих, чем детерминистские машины, можно теперь сузить до более частного вопроса: можно ли на 1-машине сделать все, что делается на p -машине? [В следующем разделе будет показано, что рассмотрение довольно узкого класса p -машин вместо более широкого класса всех мыслимых вероятностных машин не вызывает большой потери общности, что широкий класс вероятностных машин эквивалентен (в определяемом далее смысле) 1/2-машинам и что многие свойства машин широкого класса выводятся из свойств 1/2-машин.]

Необходимо еще более ограничить наш вопрос, ибо ясно, что нельзя дать никакого точного ответа, пока нет точного определения выполняемой машинами задачи.

В настоящей статье задача, выполняемая машинами, состоит в перечислении множеств. Эта задача не так уж узка, как это может показаться с первого взгляда. Например, машина может вычислять функцию, перечисляя символы « $f(n) = m$ », или устанавливать истинность или ложность некоторого высказывания о целых положительных числах перечислением символов « $P(n)$ истинно» и « $P(n)$ ложно»; или же может вычислять двоичное разложение действительного числа перечислением символов « n -й знак числа r равен a ».

Выше было определено, что понимается под 1-перечислимым множеством символов. А именно некоторое множество 1-перечислимо, если существует машина, которая, получая на входе ленту с одними только единицами, производит в итоге это множество на выходе. Необходимо дать определение перечислимости для p -машин, чтобы сравнить их с 1-машинами.

Будет предложено два определения: определение *p -перечислимых множеств* и определение *сильно p -перечислимых множеств*. Наш первоначальный неформальный вопрос будет сведен к двум точным вопросам: каждое ли p -перечислимое множество 1-перечислимо? Каждое ли сильно p -перечислимое множество 1-перечислимо?

Остается дать эти два определения. Каждому множеству возможных выходных символов данной p -машины можно приписать вероятность его появления в качестве полного выхода этой машины. Это есть вероятность того, что данное множество будет в точности множеством символов, которые появятся на выходной ленте машины, если машину привести в начальное положение, присоединить случайное устройство и пустить в работу. Указанная вероятность, конечно, равна нулю для большинства множеств. Дадим несколько примеров, ссылаясь на описанные выше машины. Предположим, что они присоединены к случайному устройству с $p = 1/2$ и стали поэтому $1/2$ -машинами. Машина № 1 перечисляет любое множество с нулевой вероятностью. Следует отметить, что конечные множества появиться не могут, и по этой причине они имеют нулевую вероятность. Хотя некоторые бесконечные множества и могут быть выходом, однако это случается лишь с нулевой вероятностью. Не следует забывать о различии между появлением с нулевой вероятностью и невозможностью появления. Машина № 2 перечисляет множество всех $[n, g(n)]$ с вероятностью 1 (и притом с достоверностью). Выход машины № 3 есть множество, состоящее из единственного символа 0, с вероятностью 1, но без достоверности. Выход машины № 5 есть множество всех целых положительных чисел — с вероятностью 1, но без достоверности, ибо бесконечная последовательность нулей и единиц содержит произвольно длинные

группы единиц с вероятностью 1, но без достоверности. Машина № 6 перечисляет с вероятностью 2^{-r} множество всех $[n, h(r, n)]$, где r фиксировано, а n пробегает все целые положительные числа. Иными словами, машина вычисляет функцию h_r [определяемую соотношением $h_r(n) = h(r, n)$] с вероятностью 2^{-r} .

О п р е д е л е н и е. Для заданного p множество символов S будет называться *сильно p -перечислимым*, если существует p -машина, производящая (в любом порядке) это множество выходных символов с положительной вероятностью.

Предыдущие примеры показывают, что следующие множества сильно $1/2$ -перечислимы: множество всех $[n, g(n)]$ — множество, состоящее только из 0; множество всех целых положительных чисел; множество всех $[n, h(r, n)]$ с фиксированным r .

Ясно, как употребить p -машину для того, чтобы дать информацию о каком-нибудь сильно p -перечислимом множестве S , производимом ею с положительной вероятностью. Заставим p -машину работать, зная, что перечисляемое ею множество будет S , исходя из вероятности появления S в качестве выхода. Поскольку S обладает положительной вероятностью, p -машину можно использовать с некоторой степенью уверенности. Если же вероятность появления S в качестве выхода равняется нулю, то нет уверенности в способности машины перечислить *в точности* S . Однако даже машина, обладающая высокой вероятностью правильно записать любой отдельный выходной символ, может перечислять S с нулевой вероятностью. Чтобы показать это, введем более слабое определение.

Существует определенная вероятность того, что данная p -машина M в *конце концов* запишет данный выходной символ. Пусть S_M множество всех выходных символов, которые машина M в конце концов запишет с вероятностью большей $1/2$. Например, присоединим рассмотренные выше машины к случайному устройству с $p = 3/4$, так что они станут $3/4$ -машинами. S_M для машины № 2 есть множество всех $[n, g(n)]$; для машины № 3 — множество, состоящее только из 0; для машины № 4 — множество, состоящее только из 1; для машины № 5 — множество целых положительных чисел и для машины № 6 — пустое множество.

Если p -машина M начала работать, то она, наконец, запишет с вероятностью большей $1/2$ любой данный элемент из S_M и с вероятностью меньшей $1/2$ любой данный элемент, не принадлежащий S_M . Отсюда ясно, как употребить машину M , чтобы получить информацию о S_M . Существует определенная степень уверенности в том, что любой элемент множества S_M в конце концов появится в качестве выходного символа и что любой элемент, не входящий в S_M , не появится в качестве выходного символа. Однако само множество S_M может при этом иметь нулевую вероятность появления, и в таком случае нет уверенности в том, что машина

перечислит в точности S_M . (Ясно, что все это остается справедливым, если в определении S_M взять вместо $1/2$ большее число. Остаются справедливыми и все последующие результаты.)

О п р е д е л е н и е. Множество символов S называется p -перечислимым, если существует такая машина M , что S есть в точности S_M .

Условие p -перечислимости множества совсем слабое, предполагается, что нет более слабого определения, которое все же предусматривало бы существование p -машины, дающей некоторую информацию о каждом элементе множества. Следует заметить, что p -машина p -перечисляет только одно множество, тогда как сильно p -перечислять она может несколько множеств или ни одного.

На первый взгляд понятие сильной p -перечислимости кажется намного сильнее понятия p -перечислимости, и действительно покажем, что сильно p -перечислимое множество p -перечислимо. Несколько более глубоким и, быть может, неожиданным является обратный результат, означающий фактическую эквивалентность этих двух понятий.

Чтобы сформулировать ключевой результат настоящей статьи и дать с его помощью ответ на два поставленных нами общих вопроса, требуется еще одно определение. Пусть p — действительное число, $0 < p < 1$. Тогда под A_p понимается бесконечная последовательность нулей и единиц $(a_1, a_2, a_3, \dots)$, где a_n есть n -й знак двоичного разложения действительного числа p , т. е. $p = 0, a_1 a_2 a_3 \dots$. Так как A_p — бесконечная последовательность нулей и единиц, то она может служить входом машин, и будет можно говорить об A_p -машинах, A_p -перечислимости и т. д.

Теперь для любого числа p между 0 и 1 у нас есть три понятия: A_p -перечислимости, p -перечислимости и сильной p -перечислимости. Из них лишь первое понятие является строго детерминистским, тогда как остальные — вероятностными.

Основной результат формулируется следующим образом.

Т е о р е м а 1. Пусть S — множество символов, p — действительное число, $0 < p < 1$. Три следующих предложения эквивалентны:

- 1) S является A_p -перечислимым;
- 2) S является p -перечислимым;
- 3) S является сильно p -перечислимым.

Доказательство этой теоремы дано в приложении, но набросок доказательства приведем здесь, показав сначала, что эта теорема дает непосредственный ответ на наши два вопроса.

Пусть p — невычислимое действительное число между 0 и 1. Известно, что такие числа существуют и что в действительности почти все (в смысле меры Лебега) числа невычислимы. Невычисли-

мость p эквивалентна невычислимости A_p . Тогда в силу леммы 3 существует множество S , которое A_p -перечислимо, но не 1-перечислимо. В силу теоремы 1 это множество как p -перечислимо, так и сильно p -перечислимо, но не 1-перечислимо. Следовательно, если имеются случайные устройства, записывающие символ 1 с вероятностью, которая не является вычислимым действительным числом, то можно построить p -машины, которые будут «перечислять» множества, не являющиеся 1-перечислимыми.

Положение совершенно меняется в случае, когда p есть вычисляемое действительное число (в частности, когда оно равно $1/2$). Это равносильно вычислимости последовательности A_p . Теорема 1 показывает, что любое p -перечислимое или сильно p -перечислимое множество должно быть A_p -перечислимым. Так как A_p -вычислимо, то по лемме 2 рассматриваемое множество должно быть 1-перечислимым. Следовательно, p -перечислимое или сильно p -перечислимое множество должно быть 1-перечислимым. Отсюда вывод, что если в качестве p брать только вычисляемые действительные числа (в частности, $1/2$), то машина со случайным устройством не сможет «перечислить» ничего сверх того, что перечисляет детерминистская машина.

Оба случая охватывает следующая теорема.

Теорема 2. Если p — вычисляемое действительное число, то любое p -перечислимое или сильно p -перечислимое множество 1-перечислимо. Если p — невычисляемое действительное число, то существуют p -перечислимые и сильно p -перечислимые множества, не являющиеся 1-перечислимыми.

Дадим набросок доказательства теоремы 1. Доказательство состоит из вывода предложения 2 из предложения 3, вывода предложения 1 из предложения 2 и вывода предложения 3 из предложения 1. Эта цепочка заключений дает требуемую эквивалентность.

То, что 3 влечет 2, доказывается тем, что с помощью p -машины, у которой выход S появляется с положительной вероятностью, строится новая машина, у которой S появляется с вероятностью большей $1/2$, а любой элемент, не входящий в S , появляется с вероятностью меньшей $1/2$, так что S перечислимо. Следовательно, всякое сильно p -перечислимое множество p -перечислимо.

То, что 2 влечет 1, доказывается следующим образом. Пусть дана p -машина M . Построим A_p -машину, A_p -перечисляющую множество S_M . Предположим, что имеется бесконечная лента, на которой записано A_p . Используя все более длинные начальные отрезки бесконечной последовательности A_p , можно вычислить сколь угодно хорошие нижние границы для вероятности справедливости высказываний вида «за время, в течение которого

из случайного устройства появилось n входных символов, M запишет символ s . Но машина M запишет s с вероятностью большей $1/2$ тогда и только тогда, когда существует такое n , что это высказывание справедливо с вероятностью большей $1/2$. Это может случиться лишь тогда, когда одна из сколь угодно хороших нижних границ (вычислимых с использованием начальных отрезков последовательности A_p) для вероятности справедливости указанного высказывания больше $1/2$. Можно вычислить *последовательно* (используя все более длинные начальные отрезки последовательности A_p) все нижние границы этих вероятностей для каждого n и каждого s . После получения нижней границы больше $1/2$ записывается соответствующий символ s . Теперь видим, что перечисленное таким образом множество символов есть S_M и что перечисление произведено посредством процесса, действительно составляющего A_p -машину. Следовательно, всякое p -перечислимое множество A_p -перечислимо.

То, что 1 влечет 3, доказывается следующим образом. Строится p -машина, вычисляющая с вероятностью больше $1/2$ двоичное разложение действительного числа p ; или выход машины есть (с вероятностью больше $1/2$) множество всех (n, a_n) , где a_n есть n -й знак в двоичном разложении числа p . Машина, делающая это, является видоизменением машины № 7. Если выход этой p -машины использовать как вход другой машины M , то выход последней будет с вероятностью больше $1/2$ в точности таким, каким был бы выход машины M , имеющий в качестве входа последовательность A_p . Если S есть любое A_p -перечислимое множество и M — машина, A_p -перечисляющая его, то построенная выше сложная машина есть p -машина, сильно p -перечисляющая S . Следовательно, любое A_p -перечислимое множество сильно p -перечислимо. Это доказывает, что 1 влечет 3, и набросок доказательства теоремы 1 тем самым закончен.

Исследуем теперь, останутся ли результаты теми же, если учитывать не только совокупность выходных символов, но и порядок их следования. Будет показано, что такое положение можно свести к частному случаю уже полученных результатов.

О п р е д е л е н и е. Бесконечная последовательность символов $S = s_1, s_2, s_3, \dots$ называется *A-вычислимой*, если она появляется в данном порядке следования на выходе некоторой A -машины. Она называется *1-вычислимой*, если A состоит целиком из единиц. S *сильно p -вычислима*, если существует p -машина, производящая S в данном порядке с положительной вероятностью. S называется *сильно p -вычислимой последовательностью*, если существует p -машина, у которой n -й выходной символ есть n -й символ в S с вероятностью больше $1/2$.

Будет показано, что для этих новых понятий справедлив точный аналог теоремы 2.

Пусть S — последовательность символов $(s_1, s_2, s_3, \dots)$. Через S' обозначим множество всех пар (n, s_n) , где n пробегает все целые положительные числа. Следующие леммы, доказываемые в приложении, суть непосредственные следствия определений.

Лемма 4. Последовательность S 1-вычислима тогда и только тогда, когда множество S' 1-перечислимо.

Лемма 5. Если S сильно r -вычислима, то S' сильно r -перечислимо. Если S r -вычислима, то S' r -перечислимо.

Из двух предыдущих лемм и теоремы 2 вытекает лемма 6.

Лемма 6. Пусть r — вычислимое действительное число (в частности, $1/2$). Тогда любая r -вычислимая или сильно r -вычислимая последовательность 1-вычислима.

Доказательство. Пусть последовательность S будет r -вычислима или сильно r -вычислима. Тогда S' r -перечислимо или сильно r -перечислимо по лемме 5. По теореме 2 множество S' будет 1-перечислимо. По лемме 4 последовательность S должна быть 1-вычислимой.

Это охватывает случай, когда r — вычислимое действительное число. Если r невычислимо, то доказательство теоремы 1, данное в приложении, показывает, что существует r -вычислимая и сильно r -вычислимая, но не 1-вычислимая последовательность (эта последовательность представляет собой двоичное разложение числа r). Объединяя оба случая, получим следующий аналог теоремы 2 для понятий вычислимости.

Теорема 3. Если r есть вычислимое действительное число, то любая r -вычислимая или сильно r -вычислимая последовательность 1-вычислима. Если r — невычислимо, то существует r -вычислимая и сильно r -вычислимая последовательность, которая не 1-вычислима.

Как частный случай этого находим, что $1/2$ -машина не может записать с положительной вероятностью никакой невычислимой последовательности.

Более общий класс машин

В этом разделе обобщается результат, даваемый теоремой 2 для вычислимого числа r . Определим широкий класс вероятностных машин — вычислимо стохастических машин (или в.с.-машин), не покрываемых нашими предыдущими определениями, и предложим

понятие эквивалентности для вероятностных машин. Покажем, что любая машина нового класса эквивалентна $1/2$ -машине, и, как следствие теоремы 2, получим, что любое множество, перечислимое машиной этого класса, должно быть 1-перечислимым. Далее, существует эффективный процесс, позволяющий по описанию машины этого класса найти описание эквивалентной ей $1/2$ -машины и описание 1-машины, 1-перечисляющей то же самое множество.

Для простоты изложения рассмотрим в этом разделе только одно определение перечислимости, соответствующее p -перечислимости, и не дадим более сильного определения, соответствующего сильной p -перечислимости. Можно показать, что те же результаты справедливы и для сильного определения.

Сначала определим класс стохастических машин. *Стохастическая машина* есть объект со счетным числом состояний $X_1, X_2, X_3, \dots$, выделенным начальным состоянием X_0 и счетным набором выходных символов $s_1, s_2, s_3, \dots$. При этом указано правило, дающее вероятность прихода машины в состояние X_p и записи ею выходного символа s_q , если перед тем она прошла последовательно через состояния $X_0, X_{j_1}, \dots, X_{j_n}$. Эту вероятность обозначим через $P(X_{j_1}, \dots, X_{j_n}; X_p; s_q)$, не требуя, чтобы это правило давало эффективный метод для вычисления вероятностей P .

Любой p -машине M можно поставить в соответствие стохастическую машину в указанном выше смысле. Состояния соответствующей стохастической машины суть конечные последовательности нулей и единиц $(a_1, \dots, a_n)$, а начальное состояние — «пустая» последовательность $(\cdot)$. (Фактически «состоянием» p -машины называют входную последовательность, воспринятую ею. Эти «состояния» могут не соответствовать возможным внутренним состояниям конкретной машины, но определяют их полностью.) Вероятность перехода из состояния $(a_1, \dots, a_n)$ в $(a_1, \dots, a_n, 1)$ есть p , а из состояния $(a_1, \dots, a_n)$ в $(a_1, \dots, a_n, 0)$ есть $1 - p$. Соответствующая стохастическая машина записывает по приходе в состояние $(a_1, \dots, a_n, a_{n+1})$ то же, что и p -машина, принявшая $(n + 1)$ -й входной символ a_{n+1} после восприятия $(a_1, \dots, a_n)$. Тем самым стохастическая машина полностью задана. Она будет в дальнейшем называться *стохастической машиной, соответствующей p -машине M* . Таким образом, p -машины можно отождествить с частными случаями стохастических машин. Легко видеть, что если p — вычислимое число, то соответствующая стохастическая машина обладает следующим специальным свойством: функция P *вычислима*, т. е. существует *эффективный процесс*, позволяющий по заданному целому положительному числу m состояниям $X_{j_1}, \dots, X_{j_n}$ и X_p и выходному символу s_q вычислить первые m знаков двоичного разложения $P(X_{j_1}, \dots, X_{j_n}; X_p; s_q)$.

Назовем стохастические машины с таким свойством вычислимо стохастическими машинами (или в.с.-машинами). Непосредственным следствием этих определений является следующая лемма.

Л е м м а 7. *Стохастическая машина, соответствующая p -машине, является в.с.-машиной тогда и только тогда, когда p — вычислимое действительное число.*

Остальную часть этого раздела посвятим тому, чтобы показать, что результат, сформулированный в теореме 2 для p -машин с вычислимым p , также справедлив и для гораздо более широкого класса в.с.-машин. Таким образом, любое множество, перечислимое в.с.-машиной, 1-перечислимо.

Как для стохастических, так и для p -машин можно говорить о вероятности того, что машина в конце концов запишет выходные символы $s_{j_1}, \dots, s_{j_n}$ (подобно тому, как раньше говорилось о вероятности того, что p -машина в конце концов запишет одиночный символ s). Поскольку наши машины служат для перечисления множеств, то любые две машины, которые в конце концов сделают одно и то же с одной и той же вероятностью, можно считать тождественными с точки зрения окончательного выхода. Поэтому можно ввести следующее определение.

О п р е д е л е н и е. Два объекта — стохастические машины или p -машины — будут называться *эквивалентными*, если для любого конечного множества символов $s_{j_1}, \dots, s_{j_n}$ они имеют одинаковую вероятность включения (в конце концов) этого множества в их выход. Если эти машины суть M_1 и M_2 , то их эквивалентность будем обозначать так: $M_1 \sim M_2$.

Например, p -машина и соответствующая ей стохастическая машина эквивалентны.

Вспомним, что p -машина названа p -перечисляющей некоторое множество символов, если каждый элемент этого множества появляется на выходе машины с вероятностью большей $1/2$, а любой выходной символ, не принадлежащий данному множеству, появляется с вероятностью не большей $1/2$. Определение можно непосредственно применить к стохастическим машинам и в.с.-машинам. Если M — стохастическая машина, то S_M будет множеством всех выходных символов, появляющихся с вероятностью большей $1/2$, в согласии с введенным выше обозначением для p -машин.

S_M есть множество символов, которое перечисляет машина M . Как непосредственное следствие из нашего определения эквивалентности получаем, что $M_1 \sim M_2$ влечет $S_{M_1} = S_{M_2}$, т. е. две эквивалентные машины перечисляют одно и то же множество.

О п р е д е л е н и е. Множество S в.с.-перечислимо, если существует такая в.с.-машина M , что $S = S_M$.

Теперь ответим на вопрос: каждое ли в.с.-перечислимое множество 1-перечислимо? Ответ вытекает из следующего утверждения, доказанного в приложении.

Теорема 4. *Каждая в.с.-машина эквивалентна $1/2$ -машине.*

Следовательно, если M есть в.с.-машина, то множество S_M , которое перечисляет машина M , перечислимо также некоторой $1/2$ -машинной. По теореме 2 множество S_M должно быть 1-перечислимым. Отсюда вытекает следующее утверждение.

Теорема 5. *Любое в.с.-перечислимое множество 1-перечислимо.*

Этот результат и был упомянут выше.

В действительности доказательство теоремы 4 дает больше, чем было высказано. Могло случиться, что доказательство только устанавливает существование $1/2$ -машины, эквивалентной произвольной в.с.-машине, но не дает эффективных средств для ее нахождения. Однако это не так, и доказательство дает эффективный процесс.

Теорема 4 (дополнение). *Существует эффективный процесс, позволяющий по заданному описанию в.с.-машины получить описание эквивалентной ей $1/2$ -машины.*

Следовательно, для заданного описания в.с.-машины M_1 существует эффективный процесс, позволяющий найти такую $1/2$ -машину M_2 , что $S_{M_1} = S_{M_2}$. Но это еще не дает эффективного расширения теоремы 5, для которого необходимо следующее эффективное расширение части теоремы 1.

Теорема 1 (дополнение). *Существует эффективный процесс, позволяющий по заданному описанию p -машины M (где p — вычислимое число) получить описание 1-машины, 1-перечисляющей то же самое S_M , которое p -перечисляет машина M .*

Из объединения дополнений теорем 1 и 4 следует теорема 5.

Теорема 5 (дополнение). *Существует эффективный процесс, позволяющий по заданному описанию в.с.-машины, в.с.-перечисляющей множество S , получить описание 1-машины, которая 1-перечисляет S .*

Результаты относительно вычислимости последовательностей, приведенные в конце предыдущего раздела, остаются справедливыми для в.с.-машин. Очевидным образом можно определить в.с.-вычислимые последовательности. Доказательства лемм 4, 5 и 6 в этом случае справедливы, откуда следует, что любая в.с.-вычислимая последовательность 1-вычислима.

Приложение

Лемма 1. *Последовательность A вычислима тогда и только тогда, когда соответствующее множество A' 1-перечислимо.*

Доказательство. Допустим, что последовательность $A = (a_1, a_2, a_3, \dots)$ вычислима. Определим машину, положив $f(b_1, \dots, b_n) = [(1, a_1), \dots, (n, a_n)]$. Ввиду вычислимости A , это соглашение действительно определяет машину. Эта машина 1-перечисляет множество A' . Обратно, допустим, что множество A' 1-перечислимо. Тогда существует 1-машина, окончательным выходом которой является множество всех символов (n, a_n) , хотя они и могут появляться в другом порядке. Чтобы определить a_n , просматривается выходная лента машины вплоть до появления пары, первый член которой есть число n . В конце концов это заведомо случится. Второй член указанной пары будет a_n . Таким образом, существует эффективный процесс для определения a_n ; следовательно, A есть вычислимая последовательность.

Лемма 2. *Если A есть вычислимая последовательность, то любое A -перечислимое множество 1-перечислимо (и обратно).*

Доказательство. Пусть машина M перечисляет множество S для вычислимой последовательности $A = (a_1, a_2, a_3, \dots)$ в качестве входа. Вследствие леммы 1 существует 1-машина, выходом которой является множество всех (n, a_n) . Этот выход можно эффективно преобразовать в последовательность $(a_1, a_2, a_3, \dots)$ и использовать как вход для M . Такое соединение есть 1-машина, 1-перечисляющая S . Справедливость обратного, т. е. того, что любое 1-перечислимое множество A -перечислимо, очевидна. В самом деле, 1-машину, 1-перечисляющую множество S , можно преобразовать в машину, выход которой не зависит от входа и которая при любом заданном входе перечисляет S .

Лемма 3. *Если A является невычислимой последовательностью, то существуют A -перечислимые множества, которые не являются 1-перечислимыми.*

Доказательство. Множество A' является таким множеством. Теперь формализуем понятие случайного устройства. Пусть D — множество всех бесконечных последовательностей $A = (a_1, a_2, a_3, \dots)$ из нулей и единиц. Пусть $C(a_1, \dots, a_n)$ — подмножество D , состоящее из всех последовательностей, которые начинаются с $(a_1, \dots, a_n)$. Будем рассматривать D как измеримое пространство [5], в котором классом измеримых множеств является σ -кольцо [5], порожденное множествами $C(a_1, \dots, a_n)$. Пусть M — машина (точное определение машины уже было дано) и S — множество возможных выходных символов машины M .

Пусть Q_S — множество всех конечных или бесконечных последовательностей $T = (s_{j_1}, s_{j_2}, \dots)$, элементы которых входят в S . Будем рассматривать Q_S как измеримое пространство, в котором классом измеримых множеств является σ -кольцо, порожденное множествами $C(s_{j_1}, \dots, s_{j_n})$ и конечными последовательностями. Пусть R_S — множество всех подмножеств в S и, кроме того, пусть $E(j_1, \dots, j_n; k_1, \dots, k_m)$ — подмножество в R_S , состоящее из всех подмножеств в S , содержащих s_{j_r} и не содержащих s_{k_r} . Рассмотрим R_S как измеримое пространство, в котором классом измеримых множеств является σ -кольцо, порожденное множествами E .

Машина M ставит в соответствие каждой входной бесконечной последовательности из D некоторую выходную последовательность из Q_S . Это определяет функцию $f_M: D \rightarrow Q_S$. С другой стороны, отображение $q_S: Q_S \rightarrow R_S$ переводит последовательность в множество ее элементов. Сложное отображение $q_S f_M$ обозначим через h_M . Если A входит в D , то $h_M(A)$ есть множество, которое перечислит машина M , снабженная входом A .

Из этих определений непосредственно следует лемма.

Л е м м а. *Отображения f_M , q_S и h_M сохраняют измеримость.*

Пусть p — действительное число, $0 < p < 1$. D можно рассматривать как пространство выборки случайного устройства, записывающего единицы с вероятностью p и нули с вероятностью $1 - p$. Иными словами, D есть пространство всех возможных событий, какие могут произойти, если наблюдать бесконечно много независимых операций случайного устройства. Каждому p естественным образом соответствует мера m_p на D , приписывающая каждому измеримому подмножеству в D вероятность того, что случайное устройство запишет бесконечную последовательность из этого подмножества¹⁾. Поэтому m_p индуцирует обычным образом вероятностные меры на пространствах Q_S и R_S . Каждому измеримому подмножеству E в Q_S приписывается вероятность $m_p[f_M^{-1}(E)]$, а каждому измеримому подмножеству F в R_S приписывается вероятность $m_p[h_M^{-1}(F)]$. Эти индуцированные меры будем обозначать через $m_{M,p}$, а иногда через m_M . Из контекста всегда будет ясно, какая мера используется.

Значение мер $m_{M,p}$ очевидно. Если E есть измеримое подмножество в Q_S , то $m_{M,p}(E)$ представляет вероятность того, что выходная последовательность машины M входит в E , когда M работает как p -машина. Если F — измеримое подмножество в R_S , то $m_{M,p}(F)$ — вероятность того, что множество выходных символов, перечисляемых M , есть элемент из F , когда машина M работает как p -машина.

¹⁾ По этому поводу см. [3].

Если U — какое-нибудь подмножество в S , то множество $\{U\}$, состоящее только из U , есть измеримое подмножество в R_S . Следовательно, всегда можно говорить о вероятности некоторого множества U быть выходом некоторой p -машины. $E(j_1, \dots, j_n)$ есть подмножество в R_S , состоящее из всех множеств, которые содержат все s_{j_r} . Ввиду измеримости $E(j_1, \dots, j_n)$ можно во всех случаях говорить о вероятности того, что p -машина в конце концов запишет некоторое конечное множество выходных символов.

Теперь дадим доказательство теоремы 1. Сначала докажем, что 3 влечет 2, т. е. что сильно p -перечислимое множество S должно быть p -перечислимым.

Пусть машина M сильно p -перечисляет множество S , т. е. ее окончательный выход есть с положительной вероятностью множество S , если машину соединить со случайным устройством, у которого вероятность записи единиц равна p . Достаточно будет найти новую машину M' , выход которой есть S с вероятностью большей $1/2$, если M' соединит с таким же случайным устройством. Машина M' при этом p -перечислит множество S , ибо каждый элемент из S появится с вероятностью большей $1/2$, а каждый элемент, не входящий в S , — с вероятностью меньшей $1/2$.

Построим машины $M(b_1, \dots, b_n)$, описываемые так: выход машины $M(b_1, \dots, b_n)$ в ответ на вход $(a_1, \dots, a_m)$ такой же, как и выход машины M в ответ на вход $(b_1, \dots, b_n, a_1, \dots, a_m)$. Следовательно, машина $M(b_1, \dots, b_n)$ действует в своем начальном состоянии так же, как если бы она была машиной M , уже воспринявшей вход $(b_1, \dots, b_n)$. Интуитивно представляется довольно правдоподобным, что если множество S есть выход с положительной вероятностью машины M , то существуют машины $M(b_1, \dots, b_n)$, выход которых есть S с вероятностью, сколь угодно близкой к 1. [Такую $M(b_1, \dots, b_n)$ можно было бы взять в качестве M' и тем самым завершить доказательство того, что 3 влечет 2.] Подтверждается это так. Пусть D_S есть $h_M^{-1}[\{S\}]$, т. е. подмножество в D , состоящее из всех последовательностей, дающих на выходе множество S , если сделать их входом машины M . Так как M производит S с положительной вероятностью, то $m_p(D_S) > 0$. Вспомним, что $C(b_1, \dots, b_n)$ есть подмножество в D , состоящее из всех последовательностей с начальным отрезком $(b_1, \dots, b_n)$. Вероятность того, что S будет выходом $M(b_1, \dots, b_n)$, равна

$$m_p[D_S C(b_1, \dots, b_n)] / m_p[C(b_1, \dots, b_n)].$$

Следовательно, наша цель будет достигнута, если удастся найти такие последовательности $(b_1, \dots, b_n)$, чтобы

$$m_p[D_S C(b_1, \dots, b_n)] / m_p[C(b_1, \dots, b_n)]$$

было произвольно близко к 1. В действительности справедливо более сильное утверждение.

Лемма. Пусть $B = (b_1, b_2, b_3, \dots)$ — некоторая последовательность в D . Тогда для всех B в D , кроме множества меры нуль,

$$\lim_{n \rightarrow \infty} m_p [D_S C(b_1, \dots, b_n)] / m_p [C(b_1, \dots, b_n)]$$

существует и равен 1 для каждой последовательности из D_S и равен 0 для каждой последовательности вне D_S .

Для доказательства этой леммы установим сохраняющее меру соответствие между пространством D с мерой m_p и единичным интервалом I с мерой Лебега и воспользуемся следующим результатом о мере Лебега m , являющимся точным аналогом нашей леммы.

Теорема о метрической плотности¹⁾. Пусть F — измеримое подмножество в I , x — точка в I и $I_n(x)$ — убывающая последовательность интервалов, пересечение которых есть x . Тогда для всех x , кроме множества меры нуль,

$$\lim_{n \rightarrow \infty} m[F \cap I_n(x)] / m[I_n(x)]$$

существует и равен 1 для x в F и равен 0 для x вне F .

Соответствие устанавливается так. Пусть отображение $f: D \rightarrow I$ задается формулой

$$f(A) = \sum_{n=1}^{\infty} a_n 2^{-n},$$

где $A = (a_1, a_2, a_3, \dots)$. Образ каждой бесконечной последовательности есть действительное число, двоичным разложением которого является эта последовательность. Отображение f есть отображение «на» и взаимнооднозначно, кроме счетного множества точек. Отображение f и мера m_p индуцируют меру m'_p на I . Если $p = 1/2$, то m'_p есть обычная мера Лебега и f — искомое соответствие. При $p \neq 1/2$ надо сделать еще шаг. Пусть $g: I \rightarrow I$ задается посредством равенства $g(x) = m'_p(I_x)$, где I_x — замкнутый интервал от 0 до x . Ясно, что g есть отображение «на», взаимнооднозначно, монотонно и непрерывно в обе стороны. Пусть $h: D \rightarrow I$ есть сложное отображение, составленное из f и g . Это h есть искомое соответствие. Его свойства таковы.

1) если E — измеримое подмножество в D , то $h(E)$ есть измеримое по Лебегу подмножество в I и $m[h(E)] = m_p(E)$;

2) если $B = (b_1, b_2, b_3, \dots)$ — последовательность в D , то $h[C(b_1, \dots, b_n)]$ есть интервал, содержащий точку $h(B)$ из I .

¹⁾ Доказательство см. [6], ч. I, стр. 190. (А также книгу И. П. Натансона «Основы теории функций вещественной переменной», 2-е изд. М. Физматгиз, 1957, стр. 229. — *Прим. перев.*)

Ясно, что так как отображение h обладает этими свойствами, то для m_p и D справедливо предложение, аналогичное теореме о метрической плотности, и что высказанная лемма содержится в этом предложении.

Таким образом, доказано, что 3 влечет 2. Доказано даже несколько больше. А именно, нами доказано, что если существует p -машина, производящая S с ненулевой вероятностью, то существуют p -машины, производящие S с вероятностью, сколь угодно близкой к 1. Заметим, что доказательство неконструктивно и не содержит эффективного процесса для нахождения таких машин.

После этого докажем, что 2 влечет 1, т. е. что p -перечислимое множество S должно быть A_p -перечислимым.

Пусть p — любое действительное число, $0 < p < 1$. Пусть M — любая p -машина. Для доказательства того, что 2 влечет 1, достаточно будет построить машину, A_p -перечисляющую S_M (множество, p -перечисляемое машиной M).

Пусть n — целое положительное число; s — возможный выходной символ машины M ; $q(n, s)$ — вероятность того, что M записала символ s за время, в течение которого были восприняты n входных символов. Заметим, что $q(n, s) \leq q(n+1, s)$ и $\lim_{n \rightarrow \infty} q(n, s)$ есть вероятность того, что символ s в конце концов будет записан. S_M есть множество всех таких символов s , что

$$\lim_{n \rightarrow \infty} q(n, s) > \frac{1}{2}.$$

Следовательно, S_M есть множество всех s , для которых существуют такие n , что $q(n, s) > 1/2$. Пусть $B_{n,s}$ — множество всех входных последовательностей $(b_1, \dots, b_n)$ длины n , которые при подаче на вход M производят выход, содержащий s . Вероятность того, что первыми n выходными символами случайного устройства будет элемент множества $B_{n,s}$ равна в точности $q(n, s)$. Пусть $q(b_1, \dots, b_n)$ — вероятность появления последовательности $(b_1, \dots, b_n)$ в качестве выхода случайного устройства. $q(b_1, \dots, b_n) = p^r(1-p)^t$, где r — число тех b , которые равны 1, а t — число тех b , которые равны 0. $q(n, s) = \sum q(b_1, \dots, b_n)$, где суммирование производится по всем последовательностям из $B_{n,s}$. Пусть действительное число p имеет двоичное разложение $0, a_1 a_2 a_3 \dots$. Пусть $p_m = 0, a_1 a_2 \dots a_m$ и $p'_m = 0, a_1 a_2 \dots a_m + 2^{-m}$, так что $p_m \leq p \leq p'_m$. Нижняя граница для $q(b_1, \dots, b_n)$ равна $r_m^* (1-p'_m)^t$. Обозначим ее через $q_m(b_1, \dots, b_n)$. Пусть $q_m(n, s) = \sum q_m(b_1, \dots, b_n)$, где суммирование производится по всем последовательностям из $B_{n,s}$; $q_m(n, s)$ есть нижняя граница для

$q(n, s)$. Так как $p_m \leq p_{m+1}$ и $p'_m \geq p'_{m+1}$, то для всех последовательностей $q_m(b_1, \dots, b_n) \leq q_{m+1}(b_1, \dots, b_n)$ и $q_m(n, s) \leq q_{m+1}(n, s)$. Так как $p_m \rightarrow p$ и $p'_m \rightarrow p$, то $q_m(b_1, \dots, b_n) \rightarrow q(b_1, \dots, b_n)$ для всех последовательностей и $q_m(n, s) \rightarrow q(n, s)$. Следовательно, S_M , т. е. множество всех s , для каждого из которых существует некоторое n с $q(n, s) > 1/2$, есть также множество всех s , для каждого из которых существуют некоторое m и некоторое n с $q_m(n, s) > 1/2$.

Теперь из построения ясно, что, имея ленту с записанной на ней последовательностью A_p , можно последовательно вычислить эффективным способом все рациональные числа $q_m(n, s)$ для всех целых положительных n и m и каждого возможного выходного символа s . Если записывать символ s каждый раз, когда встречается некоторое $q_m(n, s) > 1/2$, то перечисляемое множество будет в точности S_M и этот процесс перечисления есть A_p -машина.

Далее докажем, что 1 влечет 3, т. е. что A_p -перечислимое множество S должно быть сильно p -перечислимым.

Случай, когда p — конечная двоичная дробь, рассмотрим отдельно. В этом случае A_p вычислима. Тогда, если множество S A_p -перечисливо, то оно 1-перечисливо в силу леммы 2. Поэтому можно построить машину, которая не зависит от входа и перечисляет S при любом заданном входе. Если выход случайного устройства, обладающего вероятностью p записи 1, использовать как вход этой машины, то она перечислит S . Следовательно, S сильно p -перечисливо.

Теперь построим машину M со следующим свойством ¹⁾. Пусть p — действительное число, $0 < p < 1$, имеющее бесконечное двоичное разложение. Если выход случайного устройства, записывающего 1 с вероятностью p , подвести к M , то на выходной ленте машины M будет записано с вероятностью p двоичное разложение числа p . (Выбор числа $3/4$ несуществен. Оно может быть заменено всюду в дальнейшем любым X , $1/2 < X < 1$.)

Пусть N — машина, A_p -перечисляющая S . Если выходная лента машины M , соединенной со случайным устройством, имеющим вероятность p , служит входом в N , то выход сложной машины есть с вероятностью большей $3/4$ множество S . Сложная машина есть p -машина, сильно p -перечисляющая множество S . Следовательно, если можно построить M , то будет доказано, что любое A_p -перечислимое множество сильно p -перечисливо. Тем самым будут охвачены оба случая: когда p — конечная двоичная дробь и когда p — бесконечная двоичная дробь, так что достаточно построить такую машину M .

¹⁾ Мы обязаны Х. Троттеру указанием, приведшим к описываемому ниже построению.

Пусть $(c_1, c_2, c_3, \dots)$ — такая вычислимая последовательность рациональных чисел, что $0 < c_j < 1$ и $\prod_{j=1}^{\infty} c_j > 3/4$.

Машина M будет у нас работать следующим образом. Восприняв выход случайного устройства, наделенного вероятностью p , машина M вычисляет долю единиц после каждого выхода. Она выжидает, когда можно сказать ¹⁾ с вероятностью большей c_1 , каким должен быть первый двоичный знак числа p , чтобы дать наблюдаемую долю единиц. Тогда M записывает этот знак на выходной ленте. Затем она забывает эту наблюдаемую долю и начинает вычислять долю единиц в дальнейшем выходе случайного устройства. Она выжидает, когда можно сказать с вероятностью большей c_2 , какими должны быть первые два двоичных знака числа p , чтобы дать наблюдаемую долю единиц. Тогда она записывает на выходной ленте второй из этих знаков и начинает снова. Продолжая работать в том же роде, она записывает на своей выходной ленте правильное двоичное разложение числа p с вероятностью, большей чем $\prod_{j=1}^{\infty} c_j > 3/4$.

Точное построение таково. Пусть X_n есть n -я координатная функция на D . Это значит, что при $A = (a_1, a_2, a_3, \dots)$

$$X_n(A) = a_n.$$

Будем рассматривать D как множество всех возможных выходных последовательностей случайного устройства.

Число единиц, появившихся во время первых n операций устройства, и доля этих единиц суть соответственно

$$\sum_{j=1}^n X_j \quad \text{и} \quad Y_n = \frac{1}{n} \sum_{j=1}^n X_j.$$

Вспомним, что m_p есть мера на D , используемая, когда устройство записывает единицы с вероятностью p .

Определим функцию f следующим образом:

$$f(m, q, n, r) = m q / 2^m \{A : |Y_n(A) - q / 2^m| > r / 2^n\}$$

для всех целых положительных m, n и всех $q = 1, \dots, 2^m - 1$ и всех $r = 1, \dots, 2^n - 1$. Это вероятность того, что случайное

¹⁾ Наибольшая трудность остальной части доказательства — показать, что этой вероятности можно достичь равномерно по p , даже для p , сколь угодно близких к конечным двоичным дробям.

устройство, имеющее вероятность $p = q/2^m$, запишет последовательность n символов, в которой доля единиц отклоняется от $q/2^m$ более чем на $r/2^n$. Легко видеть, что $f(m, q, n, r)$ — рациональное число и что функция f *вычислима*. Заметим, что $f(m, q, n, 2^n) = 0$, и потому функция $g(m, q, n)$, которую мы определим как

$$\frac{1}{2^n} \left[\text{наименьшее } r, \text{ такое, что } f(m, q, n, r) < \frac{1 - c_m}{5 \cdot 2^m n^2} \right],$$

определена всегда. Так как в качестве последовательности c_m выбрана вычислимая последовательность рациональных чисел, то g есть *вычислимая* функция, принимающая только рациональные значения.

Лемма. $\lim_{n \rightarrow \infty} g(m, q, n) = 0$.

Доказательство. Теорема из работы [4] (стр. 145) показывает, что при $x = n^{1/3}/(pq)^{1/2}$

$$m_p \left\{ A: |Y_n(A) - p| > \frac{1}{n^{1/6}} \right\} \sim \sqrt{\frac{2p(1-p)}{\pi}} \frac{e^{-\frac{n^{2/3}}{2pq}}}{n^{1/3}} = o\left(\frac{1}{n^2}\right);$$

отсюда при $p = q/2^m$

$$g(m, q, n) = o\left(\frac{1}{n^{1/6}}\right)$$

и, в частности,

$$\lim_{n \rightarrow \infty} g(m, q, n) = 0.$$

Теперь работу машины можно описать так. Она воспринимает входную последовательность $A' = (a_1^1, a_2^1, a_3^1, \dots)$ из случайного устройства и вычисляет последовательно все числа $Y_1(A^1)$, $Y_2(A^1)$, $Y_3(A^1)$, $\dots$. Дойдя до n , при котором $|Y_n(A^1) - 1/2| > g(1, 1, n)$ (увидим, что это случится с вероятностью 1), она записывает на своей выходной ленте первый двоичный знак числа $Y_n(A^1)$. Покажем, что с вероятностью большей c_1 этот знак будет первым в числе p . Затем она принимает новую входную последовательность $A^2 = (a_1^2, a_2^2, a_3^2, \dots)$ из случайного устройства и вычисляет последовательно числа $Y_n(A^2)$ и числа $g(2, 1, n)$, $g(2, 2, n)$ и $g(2, 3, n)$. Дойдя до n , при котором $|Y_n(A^2) - \frac{q}{4}| > g(2, q, n)$ для $q = 1, 2, 3$, она записывает на выходной ленте второй двоичный знак числа $Y_n(A^2)$. С вероятностью большей c_2 он будет вторым двоичным знаком числа p . На m -м шаге машина работает аналогично. Она принимает входную последовательность

$A^m = (a_1^m, a_2^m, a_3^m, \dots)$, вычисляет числа $Y_n(A^m)$ и, достигнув n , при котором $|Y_n(A^m) - \frac{d}{2^m}| > g(m, q, n)$, она записывает для всех $q = 1, \dots, 2^m - 1$ на своей выходной ленте m -й знак числа $Y_n(A^m)$. С вероятностью большей c_m этот знак будет m -м знаком числа p . Таким образом, машина записывает все знаки числа p с вероятностью большей $\prod_{m=1}^{\infty} c_m > 3/4$.

Остается проверить, что m -й знак есть m -й знак числа p с вероятностью большей c_m . Определим функцию U_m на D следующим образом. Если A есть некоторая последовательность в D , то положим $U_m(A) = Y_r(A)$, где r — наименьшее n , такое, что $|Y_n(A) - \frac{q}{2^m}| > g(m, q, n)$ для всех $q = 1, \dots, 2^m - 1$, если такое целое n существует. Если такого не существует, то $U_m(A)$ считается неопределенной.

Если последовательность A служит выходом случайного устройства и используется машиной, чтобы по указанным правилам работы определить m -й знак числа p , то результат будет правилен тогда и только тогда, когда $U_m(A)$ определена и $q/2^m \leq U_m(A) < (q+1)/2^m$, где q такое, что $q/2^m \leq p < (q+1)/2^m$. (Заметим, что если $U_m(A)$ не определена, то машина не придет ни к какому результату.)

Найдем сначала вероятность того, что $U_m(A)$ не определена. $U_m(A)$ не определена тогда и только тогда, когда для каждого n существует такое q , что $|Y_n(A) - q/2^m| \leq g(m, q, n)$. По доказанной лемме

$$\lim_{n \rightarrow \infty} g(m, q, n) = 0$$

и, следовательно, существует такое единственное q_0 , что для всех достаточно больших n будет $|Y_n(A) - q_0/2^m| \leq g(m, q_0, n)$. Но тогда

$$\lim_{n \rightarrow \infty} Y_n(A) = \frac{q_0}{2^m}.$$

По усиленному закону больших чисел [4] для каждого A , кроме множества m_p -меры ноль,

$$\lim_{n \rightarrow \infty} Y_n(A) = p.$$

Так как p есть бесконечная двоичная дробь, то она не равна ни одному $q/2^m$. Следовательно,

$$m_p\{A: U_m(A) \text{ не определена}\} = 0.$$

Так как вероятность того, что U_m не определена, равна 0, то вероятность того, что первые m знаков в $U_m(A)$ правильны, равна 1 минус вероятность того, что среди первых m знаков в $U_m(A)$ есть неправильный. Поэтому достаточно вычислить последнюю вероятность. Она равна

$$\sum m_p \{A: U_m(A) \text{ определена и } q/2^m \leq U_m(A) < (q+1)/2^m\},$$

где сумма берется по всем q , для которых неверно, что $q/2^m \leq p < (q+1)/2^m$. Оценим каждое слагаемое в отдельности и покажем, что оно не превышает $(1 - c_m)/2^m$.

Отсюда будет следовать, что эта сумма не превышает $1 - c_m$ и вероятность правильности m -го разряда в $U_m(A)$ больше, чем c_m . Так для завершения доказательства осталось только сделать оценку. Возьмем любое из множеств

$$\{A: U_m(A) \text{ определена и } q/2^m \leq U_m(A) < (q+1)/2^m\},$$

для которых p не лежит между $p/2^m$ и $(q+1)/2^m$. Обозначим это множество через E . Допустим, что $p < q/2^m$. (Случай $p > (q+1)/2^m$ можно рассматривать аналогично.) В силу способа определения U_m , множество E есть сумма попарно не пересекающихся множеств $C(b_1, \dots, b_r)$, где доля единиц среди b_j заключена между $q/2^m$ и $(q+1)/2^m$. Вспомним, что $C(b_1, \dots, b_r)$ есть подмножество в D , состоящее из всех последовательностей с начальным отрезком $(b_1, \dots, b_r)$. Для каждой системы $(b_1, \dots, b_r)$ пусть B будет некоторой бесконечной последовательностью с этой системой в качестве начального отрезка. Тогда E есть сумма непесекающихся $C(b_1, \dots, b_r)$, с такими $(b_1, \dots, b_n)$, что для каждого $s < r$ существует t , $t = 1, \dots, 2^m - 1$, при котором $|Y_s(B) - t/2^m| \leq g(m, t, s)$, в то время как $|Y_r(B) - t/2^m| > g(m, t, r)$ для всех t и $q/2^m \leq U_n(B) < (q+1)/2^m$. Так как $p < q/2^m$, то $m_p[C(b_1, \dots, b_r)] < m_{q/2^m}[C(b_1, \dots, b_r)]$, если доля единиц в $(b_1, \dots, b_r)$ заключена между $q/2^m$ и $(q+1)/2^m$. Следовательно, $m_p(E) < m_{q/2^m}(E)$. Но

$$m_{q/2^m}(E) = m_{q/2^m} \{A: U_m(A) \text{ определена и } q/2^m \leq U_m(A) < (q+1)/2^m\} \leq m_{q/2^m} \{A: U_m(A) \text{ определена}\}.$$

$U_m(A)$ определена только тогда, когда существует такое n , что $|Y_n(A) - q/2^m| > g(m, q, n)$. Следовательно,

$$m_{q/2^m} \{A: U_m(A) \text{ определена}\} \leq \sum_{n=1}^{\infty} m_{\frac{q}{2^m}} \left\{ A: \left| Y_n(A) - \frac{q}{2^m} \right| > g(m, q, n) \right\} \leq \frac{1 - c_m}{5 \cdot 2^m} \sum \frac{1}{n^2} < \frac{1 - c_m}{2^m}$$

в силу определения $g(m, q, n)$. Следовательно, каждое слагаемое не превышает $(1 - c_m)/2^m$ и доказательство завершено.

Повторяем, что $3/4$ можно заменить любым X , для которого $1/2 < X < 1$, и что построение остается тем же для всех p с бесконечным двоичным разложением. Поэтому справедливо следующее. Для любого ϵ существует такая машина M_ϵ , что если использовать M_ϵ как p -машину, где p — любая бесконечная двоичная дробь, то выход машины есть двоичное разложение p с вероятностью большей $1 - \epsilon$.

Лемма 4. *S есть 1-вычислимая последовательность тогда и только тогда, когда S' есть 1-перечислимое множество.*

Доказательство. Такое же, как для леммы 1.

Лемма 5. *Если S сильно p -вычислима, то S' сильно p -перечислимо. Если S p -вычислима, то S' p -перечислимо.*

Доказательство. Пусть дана машина M . Видоизменим M так, чтобы получить машину M' , работающую следующим образом. Если выход машины M в ответ на некоторый вход есть $s_{j_1}, \dots, s_{j_n}$, то ответ машины M' на тот же вход есть $[(1, s_{j_1}) \dots (n, s_{j_n})]$. M' действительно является машиной и если M сильно p -вычисляет последовательность S , то M' сильно p -перечисляет множество S' . Если машина M p -вычисляет последовательность S , то M' p -перечисляет множество S' .

Пусть M — стохастическая машина. Пусть G — множество всех последовательностей $X = (X_{j_1}, X_{j_2}, X_{j_3} \dots)$, элементы которых суть названия состояний машины M . Пусть $C(X_{j_1}, \dots, X_{j_n})$ — множество всех последовательностей с начальным отрезком $(X_{j_1}, \dots, X_{j_n})$. Будем рассматривать G как измеримое пространство, в котором классом измеримых множеств является σ -кольцо, порожденное множествами $C(X_{j_1}, \dots, X_{j_n})$. Стохастический процесс, лежащий в основе машины, естественно, определяет меру m на G , которая сопоставляет каждому измеримому множеству вероятность прохождения машины через некоторую последовательность состояний этого множества¹⁾. Пусть S — множество возможных выходных символов машины M . Как и прежде, Q_S и R_S суть соответственно множество всех последовательностей элементов из S и множество всех подмножеств S . Отображение $f'_M : G \rightarrow Q_S$ ставит в соответствие каждой бесконечной последовательности состояний выходную последовательность, записываемую, когда машина проходит через эту последовательность состояний. Соединение

1) См. об этом [3].

этого отображения с естественным отображением $g_S: Q_S \rightarrow R_S$, переводящим последовательность во множество ее элементов, дает сложное отображение $h_M: G \rightarrow R_S$. Если X находится в G , то $h_M(X)$ есть множество символов, перечисляемых при прохождении машины через последовательность состояний X . Отображение h_M и мера m на G индуцируют меру на R_S , которую обозначим через m_M . Если E есть измеримое подмножество в R_S , то $m_M(E) = m(h_M^{-1}(E))$ есть вероятность того, что машина перечислит множество в E .

Пусть $E(j_1, \dots, j_n; k_1, \dots, k_m)$ — подмножество в R_S , состоящее из всех подмножеств в S , которые содержат $s_{j_1}, \dots, s_{j_n}$, но не содержат $s_{k_1}, \dots, s_{k_m}$ (n или m могут равняться нулю). По определению две стохастические машины M_1 и M_2 эквивалентны, $M_1 \sim M_2$, если у них одно и то же множество S возможных выходных символов и

$$m_{M_1}(E(j_1, \dots, j_n)) = m_{M_2}(E(j_1, \dots, j_n)).$$

Лемма. Если $M_1 \sim M_2$, то m_{M_1} и m_{M_2} совпадают на всех множествах вида $E(j_1, \dots, j_n; k_1, \dots, k_m)$.

Доказательство. Применим индукцию по m . Лемма верна для $m=1$ ввиду $E(j_1, \dots, j_n) = E(j_1, \dots, j_n, k_1) + E(j_1, \dots, j_n; k_1)$, где слагаемые не пересекаются. Тогда m_{M_1} и m_{M_2} совпадают на первых двух множествах и, следовательно, должны совпасть на третьем. Индуктивный переход от $m-1$ к m совершается аналогичным образом, ибо

$$\begin{aligned} E(j_1, \dots, j_n; k_1, \dots, k_{m-1}) = \\ = E(j_1, \dots, j_n, k_m; k_1, \dots, k_{m-1}) + \\ + E(j_1, \dots, j_n; k_1, \dots, k_{m-1}, k_m), \end{aligned}$$

где слагаемые не пересекаются. m_{M_1} и m_{M_2} совпадают на первых двух множествах и потому должны совпасть на третьем.

Множество всех подмножеств в R_S вида $E(j_1, \dots, j_n; k_1, \dots, k_m)$ есть кольцо множеств и порождает σ -кольцо множеств. Следовательно, если m_{M_1} и m_{M_2} совпадают на кольце, то они совпадают и на всех измеримых множествах в R_S [5, стр. 54]. Справедливость обратного очевидна, что влечет лемму.

Лемма. $M_1 \sim M_2$ тогда и только тогда, когда $m_{M_1} = m_{M_2}$.

Это показывает, что $m_{M_1} = m_{M_2}$ можно принять за определение эквивалентности вместо более слабой формулировки.

Чтобы доказать теорему 4 об эквивалентности каждой стохастической машины некоторой $1/2$ -машине, нужен еще один шаг

сведения. Стохастический процесс, лежащий в основе стохастической машины, будет сведен обычным способом к марковскому процессу.

Определение. Стохастическая машина называется марковской, если:

- 1) $P(X_{j_1}, \dots, X_{j_{n-1}}, X_{j_n}; X_p; s_q)$ не зависит от $X_{j_1}, \dots, X_{j_{n-1}}$;
- 2) для каждого X_p существует единственное s_q , для которого $P(X_{j_1}, \dots, X_{j_n}; X_p; s_q) \neq 0$.

Иными словами, не только стохастический процесс в машине является марковским, но и символ, записываемый по приходе в любое состояние, детерминированно зависит от этого и только от этого состояния.

Лемма. Каждая стохастическая машина M эквивалентна некоторой марковской стохастической машине M' . Если M есть в.с.-машина, то в качестве M' можно выбрать в.с.-машину.

Доказательство. В этом случае работает обычный прием преобразования стохастического процесса в марковский, при котором машину рассматривают как находящуюся в «состоянии» $(X_{j_1}, \dots, X_{j_n})$, если она последовательно прошла через $(X_{j_1}, \dots, X_{j_n}; s_r)$. Состояния машины M' обозначим через $X_{j_1}, \dots, X_{j_n}$. Вероятность перехода машины M' из состояния $(X_{j_1}, \dots, X_{j_n}; s_r)$ в состояние $(X_{j_1}, \dots, X_{j_n}, X_p; s_q)$ и записи символа s_q по прибытии равна $P(X_{j_1}, \dots, X_{j_n}; X_p; s_q)$, а все другие вероятности равны нулю. M' — марковская машина, и $M \sim M'$. Если M — в.с.-машина, то и M' — в.с.-машина. Кроме того, мы получили эффективный процесс, дающий по описанию M описание M' .

Таким образом, можно сосредоточить наше внимание на марковских в.с.-машинах. Это позволяет упростить обозначения. Можно считать, что состояния марковской в.с.-машины обозначены положительными числами; $P(n, m)$ будет обозначать вероятность перехода из n -го состояния в m -е. Заметим, что

$$\sum_{m=1}^{\infty} P(n, m) = 1.$$

Здесь впервые воспользуемся тем обстоятельством, что рассматриваемые объекты являются в.с.-машинами. Это означает существование эффективного процесса, позволяющего для заданных целых положительных r , n и m найти первые r знаков двоичного разложения числа $P(n, m)$. Обозначим рациональное число $0, a_1 a_2 \dots a_r$ через $P(n, m; r)$, если $(a_1, a_2, \dots, a_r)$ суть

первые r знаков двоичного разложения числа $P(n, m)$. $P(n, m; r)$ — вычисляемая функция. Пусть $Q(n, m; r)$ есть число $0, b_1 b_2 \dots b_r$, если $(b_1, b_2, \dots, b_r)$ суть первые r двоичных знаков для

$$\sum_{j=1}^m P(n, j).$$

Заметим, что $Q(n, m; r) \leq Q(n, m; r+1)$ и что

$$\lim_{r \rightarrow \infty} Q(n, m; r) = \sum_{j=1}^m P(n, j).$$

Теперь построим $1/2$ -машину, эквивалентную заданной марковской в.с.-машине. После этого теорема 4 будет доказана, так как каждая в.с.-машина эквивалентна марковской в.с.-машине. Главный прием состоит в использовании случайного устройства, записывающего единицы с вероятностью $1/2$, а также использование вычисляемой функции $Q(n, m; r)$ для получения событий, появляющихся с вероятностями $P(n, m)$. Сначала дадим упрощенный, но поучительный пример, иллюстрирующий идею построения.

Предположим, что дано случайное устройство, записывающее единицы с вероятностью $1/2$, и желательно построить устройство, записывающее единицы с вероятностью p , где p — некоторое вычисляемое число. Сначала построим 1-машину N_p , производящую двоичное разложение $0, b_1 b_2 b_3 \dots$ числа p . Пусть выход нашего случайного устройства есть последовательность $a_1, a_2, a_3, \dots$. Сравним последовательные приближения $0, b_1$; $0, b_1 b_2$; $0, b_1 b_2 b_3$; $\dots$

$\dots$; $\sum_{r=1}^n b_r 2^{-r}$; $\dots$, которые получаются для числа p при наблюдении выходной ленты машины N_p , с числами

$$0, a_1; 0, a_1 a_2; 0, a_1 a_2 a_3; \dots; \sum_{r=1}^n a_r 2^{-r}; \dots$$

Если

$$\sum_{r=1}^{\infty} a_r 2^{-r} \neq p$$

(что произойдет с вероятностью 1), то в конечном счете придем к первому n , такому, что

$$\sum_{r=1}^n a_r 2^{-r} \neq \sum_{r=1}^n b_r 2^{-r}.$$

Тогда записываем 1, если

$$\sum_{r=1}^n a_r 2^{-r} < \sum_{r=1}^n b_r 2^{-r},$$

и 0, если справедливо обратное неравенство. Так как

$$\sum_{r=1}^{\infty} a_r 2^{-r} < p$$

в точности с вероятностью p и

$$\sum_{r=1}^{\infty} a_r 2^{-r} > p$$

с вероятностью $1-p$, то событие «записана 1» появляется с вероятностью p , а событие «записан 0» — с вероятностью $1-p$. (Заметим, что событие «ничего не случилось» происходит с вероятностью 0, но все же возможно.) Повторим этот процесс. Описанный объект есть устройство, записывающее 1 с вероятностью p .

Кроме того, это построение показывает, что каждая p -машина (вспомним, что p — вычислимое число) эквивалентна некоторой $1/2$ -машине. В самом деле, любая машина M , управляемая выходом построенного выше устройства, становится $1/2$ -машиной и эквивалентна p -машине, получаемой, когда M управляется выходом случайного устройства, записывающего единицы с вероятностью p . Покажем, что конструкция $1/2$ -машины, эквивалентной заданной в. с.-машине, аналогична описанной выше.

Предположим, что дана марковская в. с.-машина M и требуется построить $1/2$ -машину M' , эквивалентную ей. Можно считать, что машина M начинает работать в состоянии 0. Пусть $(a_1, a_2, a_3, \dots)$ есть выходная последовательность случайного устройства, записывающего 1 с вероятностью $1/2$. Вероятность того, что число

$$\sum_{n=1}^{\infty} a_n 2^{-n},$$

лежит между

$$\sum_{j=1}^{m-1} P(0, j)$$

и

$$\sum_{j=1}^m P(0, j),$$

в точности равна $P(0, m)$. Пусть N есть 1-машина, вычисляющая последовательно значения вычислимой функции $Q(n, m; r)$ для каждого n, m и r . Машина M' на s -м шаге сравнивает число

$$\sum_{j=1}^s a_j 2^{-j}$$

с s первыми значениями $Q(0, m; r)$, вычисленными машиной N . С вероятностью 1 (но без достоверности) эти сравнения в конце концов дадут число t_1 , такое, что

$$\sum_{r=1}^{t_1-1} P(0, r) < \sum_{n=1}^{\infty} a_n 2^{-n} < \sum_{r=1}^{t_1} P(0, r).$$

Тогда машина M' записывает то, что машина M записала бы по приходе в состояние t_1 . Следовательно, имитируя первую смену состояний в M , машина M' записывает с вероятностью $P(0, m)$ то, что M записала бы по приходе в состояние m . Теперь M' должна имитировать вторую смену состояний в M . Она снова принимает вход $(b_1, b_2, b_3, \dots)$ из случайного устройства и сравнивает числа

$$\sum_{j=1}^s b_j 2^{-j}$$

на этот раз с s первыми значениями $Q(t_1, m; r)$, вычисленными машиной N . С вероятностью 1 эти сравнения в конце концов дадут целое число t_2 , такое, что

$$\sum_{r=1}^{t_2-1} P(t_1, r) < \sum_{n=1}^{\infty} b_n 2^{-n} < \sum_{r=1}^{t_2} P(t_1, r).$$

Тогда M' имитирует вторую смену состояний в M и записывает то, что записала бы M по приходе в состояние t_2 . Следовательно, M' записывает с вероятностью $P(t_1, m)$ то, что машина M записала бы по прибытии в состояние m . Затем M' имитирует аналогичным образом следующие смены состояний в M . Ясно, что M' есть $1/2$ -машина, эквивалентная M , а также что указанное выше построение дает эффективный процесс для преобразования описания M в описание M' .

Доказательство теоремы 4 завершено.

ЛИТЕРАТУРА

- [1] Church A., An unsolvable problem of elementary number theory, *Am. J. Math.*, 58, 345—363 (1936).
- [2] Davis M., Computability and unsolvability, New York, McGraw—Hill, 1958.
- [3] Doob J. L., Stochastic processes, John Wiley and Sons (1953). (Русский перевод: Дуб Дж. Л., Вероятностные процессы, ИЛ, М., 1956).
- [4] Феллер В., Введение в теорию вероятностей и ее приложения, ИЛ, М., 1952. Готовится к печати перевод 2 изд.
- [5] Халмош А., Теория меры, ИЛ, М., 1953.

- [6] H o b s o n E. W., The theory of functions of a real variable and the theory of Furie's series (1926).
- [7] K l e e n e S. C., Introduction to metamathematics, Van Nostrand (1952). Русский перевод: Клини С. К., Введение в метаматематику, ИЛ, М., 1957.
- [8] P o s t E. L., Finite combinatory processes — Formulation I, *J. Symbolic Logic*, 1, 103—105 (1936).
- [9] T u r i n g A. M., On computable numbers, With an application to the Entscheidungsproblem, *Proc. London Math. Soc.*, ser. 2, 43, 230—265 (1936), 43, 544—546 (1937).

БИБЛИОГРАФИЯ РАБОТ ПО ШЕННОНОВСКОЙ ТЕОРИИ ОПТИМАЛЬНОГО КОДИРОВАНИЯ ИНФОРМАЦИИ

Работы Шеннона, помещенные в данном сборнике, послужили толчком к созданию новой отрасли прикладной математики с резко выраженными границами, которая интенсивно развивается за последнее десятилетие. Эту отрасль математической науки можно назвать «шенноновской теорией оптимального кодирования информации», рассматривая ее как часть теории информации в широком смысле этого термина. Основным ее содержанием является выделение и исследование оптимальных и близких к оптимальным методов кодирования информации при передаче ее по каналам связи. Эту теорию естественно назвать шенноновской, поскольку ее основы были заложены в известной работе «Математическая теория связи»¹⁾ и, кроме того, большая часть ее идей была или впервые выдвинута, или существенно развита в работах Шеннона.

В рамках данного сборника не представляется возможным дать достаточно полный обзор истории развития шенноновской теории оптимального кодирования информации. Ввиду этого в библиографию, составленную М. А. Ратнер, включены наиболее интересные работы, имеющие существенный теоретический интерес.

Все отобранные работы разделены на 11 разделов, обозначенных буквами русского алфавита. В алфавитном списке литературы эти буквы, указывающие соответствующий раздел, стоят в конце наименования работы.

Ниже приводится содержание каждого из 11 разделов и наиболее интересные работы по каждому из них.

А. *Книги и обзорные статьи, посвященные общему изложению теории.* В качестве первого введения в теорию информации для инженеров в настоящее время можно рекомендовать книги Fano [2] и Reza [1], а также статьи Шеннона [1,3]; для математика более подходит книга Файнштейна [1]. Физические аспекты теории информации отражены в книге Бриллюэна [1]. Лучшей из популярных книг является книга Яглома и Яглома [1]. Для ознакомления с самыми современными актуальными проблемами теории можно рекомендовать книгу Wolfowitz [13] и обзорную статью Добрушина [3]. Имеется библиография Stumpers [1—4]. Кроме того, выделим следующие работы: Голдман [1], Добрушин [9], Колмогоров [1], Bell [1], Cherry [1], Hancock [1], Jakobs [1], Zemanek [1].

¹⁾ См. стр. 243 данного сборника.

Б. *Общие свойства энтропии и информации*. Основные свойства энтропии в случае дискретных величин подробно изложены в книге Файнштейна [1] и статье Хинчина [1]. Свойства информации в случае представления информации в непрерывной форме изложены во вводных разделах статьи Добрушина [5] и книги Пинскера [7]. Из работ, не отраженных в этих книгах и обзорах, отметим следующие статьи: Добрушин [8], Birch [1], Csiszár [1], Rényi [1—3], Rényi, Balatoni [1]. К этой теме относятся также следующие работы: Гельфанд, Колмогоров, Яглом [1, 2], Гельфанд, Яглом [1], Колмогоров [1], Розенблат-Рот [1], Рохлин [1], Синай [1, 2], Фадеев [1], Хинчин [1], Цзян Цзэ-пей [1], Ху Го-дин [2], Billingsley [1], Ikeda [1—3], Pérez [1], Segal [1], Tverberg [1].

В. *Исследование условий, при которых верна основная теорема Шеннона о возможности передачи информации*. Итоги первого этапа исследований, посвященного обоснованию теоремы Шеннона для дискретных каналов, подведены в монографиях Файнштейна [1] и Wolfowitz [13]. Наиболее общие результаты для непрерывных каналов содержатся в работах Добрушина [5] и Пинскера [6, 7]. Теорема Шеннона для сообщений с заданным критерием точности изучена математически строго в работах Шеннона [13] и Добрушина [8].

Отметим, кроме того, следующие работы: Блекуэлл, Брейман, Томасян [1], Перез [1], Розенблат-Рот [1—3], Хинчин [1, 2], Ху Го-дин [2], Adler [1], Blackwell [1], Breiman [1, 2], Carleson [1], Chung [1], Feinstein [1, 3], Jakobs [1, 3, 4], Kinney [1], McMillan [1], Moy Shu-Teh [1, 2], Nedoma [1, 2], Parthasarathy [1], Pérez [1, 2, 4], Takano [1], Wolfowitz [2, 4, 7, 10, 12].

Г. *Вычисление скорости передачи информации, пропускной способности и ϵ -энтропии (скорости передачи как функции искажения)*. Обзорные изложения этого вопроса до сих пор отсутствуют, хотя некоторые факты собраны в книгах Reza [1] и Faпо [2]. Вычислению пропускной способности каналов без памяти посвящены работы Шеннона [7] и Muroga [1—4]. По поводу вычисления энтропии в случае дискретного канала см. Шеннон [13], Ерохин [1]. Большая литература посвящена случаю гауссовского канала: Гельфанд, Яглом [1], Пинскер [1—6], Good, Doog [1, 2], Powers [1], Korezlioglu [1]. Существует обширная литература по вычислению пропускной способности для моделей реальных каналов связи: Gilbert [4], Добрушин, Хургин, Цыбаков [1], Линь Хай-цюань [1, 2], Овсеевич, Пинскер [1—8], Хургин [1], Сифоров [6—8], Цыбаков [1—4, 6, 7]. Отметим работы по статистической оценке энтропии и пропускной способности: Башарин [1], Добрушин [2], Lomnizky, Zaremba [1].

Выделим также следующие работы: Блекуэлл [1], Добрушин [6], Колмогоров, Тихомиров [1], Любич [1], Рубинштейн, Урбаник [1], Синай [1], Шеннон [5, 10, 11], Abramson [3], Blachman [3, 4], Chang [2], Huang, Johnson [1], Swerling [1].

Д. *Обобщения основной постановки задачи: двусторонние каналы с обратной связью; каналы с неизвестными параметрами (совместные каналы).* Некоторые результаты по этому поводу собраны в книге Wolfowitz [13]. Более общие дальнейшие результаты содержатся в работах: Добрушин [12, 14, 15], Шеннон [14], Blackwell, Breiman, Thomasian [3], Kesten [1]. Выделим также следующие работы: Добрушин [1], Шеннон [9], Bellman, Kalada [1, 2], Blackwell, Breiman, Thomasian [1], Chang [1, 3], Kelly [1, 2], Wolfowitz [8—10, 12].

Е. *Исследование вероятности ошибки при передаче новыми кодами.* В качестве введения в эту тематику рекомендуется последняя глава книги Файнштейна [1]. Наиболее сильные результаты для случая дискретного канала содержатся в последних главах книги Fano [2], а для случая непрерывного канала в работах: Шеннон [12], Thomasian [3]; см. также следующие работы: Добрушин [7, 10—13], Шеннон [7, 10, 11], Элайес [1, 2], Юшкевич [1], Abramson [3], Blackwell, Breiman, Thomasian [2], Elias [1—3], Feinstein [2], Kotz [1], Thomasian [1], Weiss [1].

Ж. *Комбинаторно-алгебраическая теория кодов, исправляющих ошибки.* Основные результаты, содержащиеся в обширной литературе по этому вопросу, собраны в книге Peterson [4], где даны также подробные ссылки на литературу.

З. *Последовательное декодирование.* Литература по этому вопросу немногочисленна; см. следующие работы: Возенкрафт, Рейффен [1], Рейффен [1], Epstein [1], Perry, Wozencraft [1], Wozencraft [1,2], Ziv [1]. Подробнее всего метод последовательного декодирования изложен в книге Возенкрафт, Рейффен [1].

И. *Статистическое кодирование сообщений (неравномерные коды).* Некоторые результаты собраны в статье: Гилберт, Мур [1]. Интересный более общий подход применен в работах: Глебский [1, 2], Левенштейн [2, 4, 5], Марков [1, 4]; см. также работы: Блох [2, 3], Мак-Миллан [1], Сардинас, Патерсон [1], Хаффмен [2], Blachman [2], Jaynes [2], Кагр [1], Neumann [2], Schützenberger, Marcus [1].

К. *Применение идей шенноновской теории оптимального кодирования информации в математической статистике.* Эта тематика подробно изложена в книге: Kullback [3]; см. также следующие работы: Айвазян [1], Линдли [1], Сакагучи [1], Hajek [1], Ikeda [2], Kullback [1, 2], Lindley [1, 2], Mourier [1], Sakaguchy [1, 2], Stam [1], Vincze [1], Watanabe [4].

Л. *Приложения шенноновской теории оптимального кодирования информации в технике связи, физике, лингвистике и т. д.* Обзор возможностей технических приложений см. Элайес [3]; по поводу приложения к физике — Бриллюэн [1]; приложения к лингвистике — Добрушин, Яглом, Яглом [1], Шеннон [16]; приложения к психологии — Quastler [1].

I. Работы на русском языке

Айвазян С. А.

1. Сравнение оптимальных свойств критериев Неймана — Пирсона и Вальда, *Теория вероятн. и ее примен.* (1959), 4, 86—93. К.

Акиндинов В. В.

1. О геометрической интерпретации двоичного кода, исправляющего кратные и одиночные ошибки, *Изв. высш. учебн. заведений, Радиофизика* (1960), 3, 4, 716—721. Ж.

Амбарцумян Г. А.

1. К энтропии цепей Маркова, *Изв. АН Арм. ССР, Сер. физ.-матем. н.* (1958), 11, 2, 31—40. Г.

Байковский В. М.

1. О пропускной способности канала с разностно-дискретной модуляцией, *Радиотехн. и электроника* (1961), 16, 7, 11—19. Г.

Бакут Л. А.

1. К теории корректирующих кодов с произвольным основанием, *Научн. докл. высш. школы, Радиотехн. и электроника* (1959), 1, 26—36. Ж.

Барнард Г. А.

1. Простые доказательства простых случаев теоремы кодирования, сб. *Теория передачи сообщений*, М., ИЛ, 1957, стр. 32—40. В.
2. Элементарное доказательство теоремы кодирования, сб. *Теория передачи сообщений*, М., ИЛ, 1957, стр. 40—43. В.

Бахметьев М. М.

1. Оценка точности определения энтропии при аппроксимации распределения вероятностей геометрической прогрессией, *Радиотехн. и электроника* (1957), 2, 6, 811—813. Г.

Башарин Г. П.

1. О статистической оценке энтропии последовательности независимых случайных величин, *Теория вероятн. и ее примен.* (1959), 4, 361—364. Г.

Блекуэлл Д.

1. Энтропия функций на цепи Маркова с конечным числом состояний, сб. *Математика* (1959), 3 : 5, 143—150. Г.

Блекуэлл Д., Брейман Л., Томасян А.

1. Доказательство теоремы Шеннона о передаче информации для неразложимых каналов с конечным числом состояний, сб. *Математика* (1960), 4 : 5, 123—135. В.

Блох Э. Л.

1. К вопросу о зависимости между скоростью передачи сообщений и помехоустойчивостью системы связи, *Радиотехн.* (1957), 12, 6, 3—15. В.
2. О передаче неравномерной бинарной последовательности равномерным кодом, *Электросвязь* (1959), 1, 76—77. И.

3. Об одном неравенстве теории информации, *Изв. АН СССР, Отд. техн. н., Энерг. и автоматика* (1961), 6, 93—100. И.

Блох Э. Л., Харкевич А. А.

1. К вопросу о геометрическом доказательстве теоремы Шеннона, *Радиотехн.* (1956), 11, 11, 5—16. В.
2. Кодирование, устойчивое по отношению к замиранию (антифединговое кодирование), *Электросвязь* (1960), 4, 3—6. Л.

Болсер М., Сильвермэн Р. А.

1. Кодирование при передаче данных с постоянной скоростью. II. Коды с исправлением многократных ошибок, сб. Коды с обнаружением и исправлением ошибок, М., ИЛ, 1956, стр. 43—59. Ж.

Бородин Л. Ф.

1. К теории корректирующих кодов с простым основанием, *Научн. докл. высш. школы, Радиотехн. и электроника* (1958), 1, 22—26. Ж.
2. Об одном регулярном способе построения корректирующих кодов, *Научн. докл. высш. школы, Радиотехн. и электроника* (1958), 1, 54—57. Ж.
3. Некоторые вопросы теории построения корректирующих кодов, сб. тр. Научно-техн. об-ва радиотехн. и электросвязи им. А. С. Попова (1958), 2, 110—151. Ж.
4. Эквидистантные и другие оптимальные и близкие к оптимальным коды, *Радиотехн. и электроника* (1960), 5, 6, 883—894. Ж.
5. Некоторые вопросы теории групповых кодов, *Радиотехн.* (1962), 8, 1275—1284. Ж.

Бородин Л. Ф., Грушко И. И.

1. О целесообразности введения интервала стирания, *Радиотехника* (1962), 17, 3, 37—47. Е.

Боуз Р. С., Рой-Чоудхури Д. К.

1. Об одном классе двоичных групповых кодов с исправлением ошибок, *Кибернетический сб.*, вып. 2, М., ИЛ, 1961, стр. 83—94. Ж.
2. Дальнейшие результаты относительно двоичных групповых кодов с исправлением ошибок, *Кибернетический сб.*, вып. 6, М., ИЛ, 1963, стр. 8—19. Ж.

Бриллюэн Л.

1. Наука и теория информации, М., Физматгиз, 1960. А.

Васильев А. М.

1. Геометрический вывод формулы для пропускной способности канала связи с шумами со специальным приемником, *Радиотехника* (1956), 2, 77—79. Г.
2. О негрупповых плотно упакованных кодах, *Проблемы кибернетики*, вып. 8, М., Физматгиз, 1962, стр. 337—339. Ж.

Варшавер Б. А.

1. К теории пропускной способности при передаче сигналов со многими дискретными значениями, *Научн. докл. высш. школы, Радиотехн. и электроника* (1958), 1, 46—50. Г.

2. К теории пропускной способности при бинарной передаче, *Радиотехн.* (1958), 13, 1, 11—21. Г.
3. К теории передачи сигналов со многими дискретными значениями, *Радиотехн.* (1959), 14, 1, 3—13. Г.
4. К сравнению равномерных кодов при бинарной передаче, сб. тр. Научно-техн. об-ва радиотехн. и электросвязи им. А. С. Попова (1959), 3, 49—58. Ж.

В а р ш а м о в Р. Р.

1. Оценка числа сигналов в кодах с корреляцией ошибок, *Докл. АН СССР* (1957), 117, 5, 739—741. Ж.
2. О методе линейного кодирования с исправлением ошибок в передаваемых сигналах, сб. тр. Научно-техн. об-ва радиотехн. и электросвязи им. А. С. Попова (1959), 3, 43—48. Ж.

В о з е н к р а ф т Дж. М., Р е й ф ф е н Б.

1. Последовательное декодирование, М., ИЛ, 1963. З.

Г а р м а ш В. А.

1. О кодировании длинных отрезков двоичных символов, *Радиотехн.* (1959), 14, 4, 62—64. Ж.
2. Способ построения оптимального двоичного кода, *Электросвязь* (1959), 2, 3—14. Ж.

Г а р м а ш В. А., К и р и л л о в Н. Е.

1. Кодирование последовательности из двух независимых символов, *Электросвязь* (1958), 9, 3—6. Ж.
2. Экспериментальное исследование статистики фототелеграфных сообщений, *Научн. докл. высш. школы, Радиотехн. и электроника* (1959), 1, 37—42. Л.

Г а р м а ш В. А., П ч е л ь н и к Б. М., К а ч е р о в и ч Я. М.

1. О кодировании сообщений равномерным статистическим кодом, *Тр. учебн. инст. связи Мин. связи СССР* (1960), 1, 17—24. Ж.

Г е л ь ф а н д И. М., К о л м о г о р о в А. Н., Я г л о м А. М.

1. К общему определению количества информации, *Докл. АН СССР* (1956), 111, 4, 745—748. Б.
2. Количество информации и энтропии для непрерывных распределений, Тр. 3-го Всес. матем. съезда, Изд. АН СССР (1958), 3, 300—320. Б.

Г е л ь ф а н д И. М., Я г л о м А. М.

1. О вычислении количества информации о случайной функции, содержащейся в другой такой функции, *Успехи матем. наук* (1957), 12, 1, 3—52, Б, Г.

Г и л б е р т Э. Н.

1. Синхронизация двоичных сообщений, *Кибернетический сб.*, вып. 5, М., ИЛ, 1962, стр. 42—59. Л.

Г и л б е р т Э. Н., М у р Э. Ф.

1. Двоичные кодовые системы переменной длины, *Кибернетический сб.*, вып. 3, М., ИЛ, 1961, стр. 103—141. И.

Глебский Ю. В.

1. Кодирование с помощью конечных автоматов, *Докл. АН СССР* (1961), 141, 6, 1054—1057. И.
2. Кодирование с помощью автоматов с конечной внутренней памятью, *Проблемы кибернетики*, вып. 8, М., 1962, стр. 127—150. И.

Голдман С.

1. Теория информации, М., ИЛ, 1957. А.

Голомб С., Велч Л. Р., Дельбрюк М.

1. Строение и свойства кодов без запятой, сб. *Математика* (1960) 4 : 5, 137—160. Ж.

Голомб С., Гордон Б., Велч Л. Р.

1. Коды без запятой, *Кибернетический сб.*, вып. 5, М., ИЛ, 1962, стр. 33—41. Ж.

Горенштейн Д., Питерсон У., Цирлер Н.

1. Квазисовершенство кодов Боуза — Чоудхури с исправлением двух ошибок, *Кибернетический сб.*, вып. 6, М., ИЛ, 1963, 20—24. Ж.

Демидович Н. Б.

1. К теории групповых кодов, *Проблемы кибернетики*, вып. 5, М., Физматгиз, 1962, стр. 105—123. Ж.

Джоши Д. Д.

1. О верхних границах для кодов с минимальным расстоянием, *Кибернетический сб.*, вып. 1, М., ИЛ, 1960, стр. 227—233. Ж.

Добрушин Р. Л.

1. Передача информации по каналу с обратной связью, *Теория вероятн. и ее примен.* (1958), 3, 4, 395—412. Д.
2. Упрощенный метод экспериментальной оценки энтропии стационарной последовательности, *Теория вероятн. и ее примен.* (1958), 3, 4, 462—464. Г.
3. Теория передачи информации, Моск. дом. науч.-техн. пропаганды им. Ф. Э. Дзержинского, Конференция по теории информации, 1958, 1. А.
4. Общая формулировка основной теоремы Шеннона в теории информации, *Докл. АН СССР* (1959), 126, 3, 474—477. В.
5. Общая формулировка основной теоремы Шеннона в теории информации, *Успехи матем. наук* (1959), 14, 6, 3—104. Б, В.
6. Оптимальная передача информации по каналу с неизвестными параметрами, *Радиотехн. и электроника* (1959), 4, 1951—1956. Д.
7. Асимптотика вероятностей ошибок при передаче информации по каналу без памяти с симметрической матрицей вероятностей перехода, *Докл. АН СССР* (1960), 133, 2, 265—268. Е.
8. Предельный переход под знаком информации и энтропии, *Теория вероятн. и ее примен.* (1960), 5, 1, 29—37. Б.
9. Математические вопросы шенноновской теории оптимального кодирования информации, сб. *Проблемы передачи информации*, АН СССР (1961), 10, 63—107. А.

10. Оптимальные бинарные коды для малых скоростей передачи информации, *Теория вероятн. и ее примен.* (1962), 7, 2, 203—213. Е, Ж.
11. Асимптотические оценки вероятности ошибки при передаче сообщения по дискретному каналу связи без памяти с симметрической матрицей вероятностей перехода, *Теория вероятн. и ее примен.* (1962), 7, 3, 283—311. Е.
12. Асимптотическая оценка вероятности ошибки при передаче сообщения по каналу без памяти с использованием обратной связи, *Проблемы кибернетики*, вып. 3, М., Физматгиз, 1962, стр. 161—168. Д, Е.
13. Асимптотическая оптимальность групповых и систематических кодов для некоторых каналов, *Теория вероятн. и ее примен.* (1963), 8, 1, 52—66. Е, Ж.
14. Единые способы передачи информации для дискретных каналов без памяти и сообщений с независимыми компонентами, *Докл. АН СССР* (1963), 148, 6, 1245—1248. Д.
15. Единые способы передачи информации; общий случай, *Докл. АН СССР* (1963), 149, 1, 16—19. Д.

Добрушин Р. Л., Хургин Я. И., Цыбаков Б. С.

1. Приближенное вычисление пропускной способности радиоканалов со случайными параметрами, *Труды Всес. совещ. по теории вероятн. и матем. стат.*, Ереван, 19—25 сент. 1958, Изд. АН Арм. ССР, 1960, 164—171. Г.

Добрушин Р. Л., Яглом А. М., Яглом И. М.

1. Теория информации и лингвистика, *Вопросы языкознания* (1960), 1, 100—110. Л.

Долуханов М. П.

1. О некоторых возможных способах измерения энтропии дискретных и непрерывных источников сообщений, *Радиотехн. и электроника* (1959), 4, 4, 559—565. Б.

Ерохин В.

1. ϵ -энтропия дискретного случайного объекта, *Теория вероятн. и ее примен.* (1958), 3, 1, 103—107. Г.

Заремба С. К.

1. Замечание к основной теореме для дискретного канала с шумами, сб. *Теория передачи сообщений*, М., ИЛ, 1957, стр. 20—28. В.

Зверев В. А., Орлов Е. Ф.

1. О скорости передачи информации по каналам с многолучевым пространством, *Изв. высш. учебн. заведений, Радиофизика* (1961), 4, 2, 282—292. Г.

Зотова Е. Н.

1. Сравнительный анализ двоичных кодов с коррекцией ошибок, *Научн. докл. высш. школы, Радиотехн. и электроника* (1958), 1, 37—45. Е, Ж.

Игнатьев Н. К.

1. О некоторых геометрических свойствах оптимального кода, *Электросвязь* (1957), 6, 3—10. Ж.

Кириллов Н. Е.

1. Коды с разделительными знаками, *Электросвязь* (1959), 6, 5—10. Ж.
2. О кодах, ограничивающих область ошибки, *Изв. АН СССР. Отд. техн. н., Энерг. и автоматика* (1959), 3, 183—184. Ж.

Кислицын С. С.

1. Средняя длина двоичного кода с минимальной избыточностью в случае, когда вероятности кодируемых символов близки друг к другу, *Теор. вероятн. и ее примен.* (1962), 7, 3, 342—343. И.

Клячкин Л. З.

1. Пропускная способность бинарной кодово-импульсной системы при неодинаковых вероятностях искажения символов, *Радиотехн.* (1958), 13, 4, 26—29. Г.

Колмогоров А. Н.

1. Теория передачи информации, сб. Сессия АН СССР по научн. пробл. автоматич. произв., 1956, Пленарн. заседания, М., АН СССР, 1957, 66—99, Дискуссия, 148—161. А.
2. Об энтропии на единицу времени как метрическом инварианте автоморфизмов, *Докл. АН СССР* (1959), 124, 4, 754—755. Б.

Колмогоров А. Н., Тихомиров В. Н.

1. ε -энтропия и ε -емкость множеств в метрических пространствах, *Успехи матем. наук* (1959), 14, 2, 3—86. Г.

Лебедев Д. С.

1. Статистические свойства множеств сообщений, *Радиотехника* (1958), 13, 1, 3—10. В.

Левенштейн В. И.

1. Об одном классе систематических кодов, *Докл. АН СССР* (1960), 131, 5, 1011—1014. Ж.
2. О некоторых свойствах кодовых систем, *Докл. АН СССР* (1961), 140, 6, 1274—1277. И.
3. Применение матриц Адамара к одной задаче кодирования, *Проблемы кибернетики*, вып. 5, М., Физматгиз, 1961, стр. 123—137. Ж.
4. Самонастраивающиеся автоматы для декодирования сообщений, *Докл. АН СССР* (1961), 141, 6, 1320—1323. И.
5. Об обращении конечных автоматов, *Докл. АН СССР* (1962), 147, 6, 1300—1303. И.

Левин Б. Р., Розанов В. С.

1. Исследование пропускной способности многоканальных систем при учете статистической структуры источника, Тр. VI Всес. совещ. по теории вероятн. и матем. статистике, 1960, Вильнюс, Гос. изд-во полит. и научн. лит. Лит. ССР, 1962, 215—222. Г.

Леммель А. Э.

1. Об одном классе кодов и их физической реализации, сб. Теория передачи сообщений, М., ИЛ, 1957, стр. 43—52. Ж.

Линдли Д. В.

1. О мере информации, даваемой экспериментом, сб. *Математика*, (1959), 3: 3, 87—104. К.

Линковский Г. Б.

1. Об информации непрерывных сообщений по Шеннону, *Научн. докл. высш. школы, Радиотехн. и электроника* (1958), 1, 12—15. Б.
2. Оценка энтропии одномерного распределения, представленного несколькими эмпирическими моментами, *Научн. докл. высш. школы, Радиотехн. и электроника* (1958), 2, 3—7. Г.

Линь Хай-цунь

1. К вычислению пропускной способности двухлучевых каналов связи, *Радиотехн. и электроника* (1961), 6, 2, 197—199. Г.
2. О пропускной способности однолучевых и двухлучевых каналов связи при поляризационных замираниях, *Радиотехн. и электроника* (1961), 6, 2, 200—203. Г.

Ллойд С. П.

1. Бинарное блочное кодирование, Кибернетический сб., вып. 1, М., ИЛ, 1960, стр. 206—226. Ж.

Любич Ю. И.

1. Замечание о пропускной способности дискретного канала связи без шумов, *Успехи матем. наук* (1962), 17, 1, 191—198. Г.

Мак-Миллан Б.

1. Два неравенства, обусловленные однозначностью расшифровывания, Кибернетический сб., вып. 3, М., ИЛ, 1961, стр. 88—92. И.

Мандельброт Б.

1. О рекуррентном кодировании, ограничивающем влияние помех, сб. Теория передачи сообщений, М., ИЛ, 1957, 139—158. И.

Марков А. А.

1. Об алфавитном кодировании, *Докл. АН СССР* (1960), 132, 3, 521—523. И.
2. Об алфавитном кодировании, *Докл. АН СССР* (1961), 139, 3, 560—561. И.
3. Нерекуррентное кодирование, Проблемы кибернетики, вып. 8, М., Физматгиз, 1962, стр. 169—186. И.
4. Условие полноты для неравномерных кодов, Проблемы кибернетики, вып. 9, М., Физматгиз, 1963, стр. 327—332. И.

Матюхин Н. Я.

1. Линейные преобразования двоичных кодов, *Автомат. и телемеханика* (1958), 19, 8, 776—787. Ж.

Мешковский К. А.

1. Оптимальные и близкие к ним двоичные коды, *Электросвязь* (1958), 5, 5—15. Ж.

2. Некоторые вопросы теории кодирования, сб. Проблемы передачи информации, АН СССР (1959), 2, 57—64. Ж.

Немировский А. С.

1. О пропускной способности многолучевого канала при разнесенном приеме с автовыбором, *Радиотехн.* (1961), 16, 9, 34—38. Г.

Новик Д. А.

1. К передаче электрических сигналов оптимальным кодом Шеннона — Фано, *Электросвязь* (1958), 7, 3—5. И.
2. Методы декодирования неравномерных двоичных кодов, *Электросвязь* (1959), 5, 3—12. И.

Овсеевич И. А., Пинскер М. С.

1. Оценка пропускной способности канала связи, параметры которого являются случайными функциями времени, *Радиотехн.* (1957), 12, 10, 40—46. Г.
2. Оценка пропускной способности некоторых реальных каналов связи, *Радиотехн.* (1958), 13, 4, 15—25. Г.
3. О пропускной способности многопутевой системы информации, *Изв. АН СССР, Отд. техн. н., Энерг. и автоматика* (1959), 1, 133—135. Г.
4. Скорость передачи информации, пропускная способность многопутевой системы и прием по методу линейно-операторного преобразования, *Радиотехн.* (1959), 14, 3, 9—21. Г.
5. Оптимальное линейное предсказание и корректирование сигнала при передаче его по многопутевой системе, *Изв. АН СССР, Отд. техн. н., Энерг. и автоматика* (1959), 2, 49—59. Г.
6. Предсказание и корректирование в канале с замираниями, *Изв. АН СССР, Отд. техн. н., Энерг. и автоматика* (1960), 3, 145—156. Г.
7. Пропускная способность каналов с общим и селективным замиранием, *Радиотехн.* (1960), 15, 12, 3—9. Г.
8. О пропускной способности многопутевой системы, *Изв. АН СССР, Отд. техн. н., Энерг. и автоматика* (1961), 4, 208—210. Г.

Остиану В. М.

1. Линейное кодирование для небинарных кодов с коррекцией ошибок, *Mathematica (RPR)* (1960), 2, 2, 291—298. Ж.
2. Построение небинарных самокорректирующихся кодов и оценка числа сигналов в них, *Докл. АН СССР* (1960), 135, 6, 1382—1384. Ж.
3. Математический метод построения небинарных кодов с коррекцией ошибок, 2^{ème} congr. math. hongrois, Budapest, 1960, Budapest, 1961, 37—41. Ж.
4. Построение небинарных кодов, корректирующих ошибки, и оценка числа сигналов в них, сб. Проблемы передачи информации, АН СССР (1961), 10, 42—48. Ж.

Перез А.

1. Теория информации с абстрактным алфавитом. Обобщенные виды предельной теоремы Мак-Миллана для случая дискретного и непре-

равного времени, *Теория вероятн. и ее примен.* (1959), 4, 1, 105—109. В.

Пинскер М. С.

1. Количество информации о гауссовском случайном процессе, содержащейся во втором процессе, стационарно с ним связанном, *Докл. АН СССР* (1954), 99, 2, 213—216. Г.
2. Вычисление скорости создания сообщений стационарным случайным процессом и пропускной способности стационарного канала, *Докл. АН СССР* (1956), 111, 4, 753—756. Г.
3. Количество информации об одном стационарном случайном процессе, содержащееся в другом стационарном случайном процессе, Тр. 3-го Всес. матем. съезда, 1, М., АН СССР, 1956, 125. Г.
4. Экстраполирование однородных случайных полей и количество информации о гауссовском случайном поле, содержащейся в другом гауссовском случайном поле, *Докл. АН СССР* (1957), 112, 5, 815—818. Г.
5. Экстраполирование случайных векторных процессов и количество информации, содержащееся в одном векторном стационарном случайном процессе относительно другого, стационарно с ним связанного, *Докл. АН СССР* (1958), 125, 1, 49—51. Г.
6. Информационная устойчивость гауссовских случайных величин и процессов, *Докл. АН СССР* (1960), 133, 1, 28—30. В.
7. Информация и информационная устойчивость случайных величин и процессов, М., Изд-во АН СССР, 1960. Б, В.
8. Энтропия, скорость создания энтропии и энтропийная устойчивость гауссовских случайных величин и процессов, *Докл. АН СССР* (1960), 133, 3, 531—534. В.

Питерсон У.

1. Кодирование и исправление ошибок для кодов Боуза — Чоудхури, Кибернетический сб., вып. 6, М., ИЛ, 1963, стр. 25—54. Ж.

Питерсон В. В., Рабин М. О.

1. О кодах для контроля логических операций, Кибернетический сб., вып. 4, М., ИЛ, 1962, стр. 105—119. Ж.

Плоткин М.

1. Двоичные коды с заданным минимальным расстоянием, Кибернетический сб., вып. 7, М., ИЛ, 1963, стр. 60—73. Ж.

Пушной Б. М.

1. Геометрическое построение оптимальных кодов, *Электросвязь* (1959), 10, 3—12. Ж.

Радченко А. Н., Мирончиков Е. Т.

1. Многотактные методы исправления одиночных и многократных близко расположенных ошибок в групповых кодах, *Радиотехн. и Электроника* (1961), 11, 1805—1812. Ж.

Рейффен Б.

1. Последовательное декодирование для каналов без памяти с дискрет-

ным входом, в кн. Возенкрафт Дж. М., Рейффен Б., Последовательное декодирование, М., ИЛ, 1963, 3.

Рид И. С.

1. Класс кодов с исправлением нескольких ошибок и схема декодирования, Кибернетический сб., вып. 1, М., ИЛ, 1960, стр. 189—205. Ж.

Рид И., Соломон Г.

1. Полиномиальные коды над некоторыми конечными полями, Кибернетический сб., вып. 7, М., ИЛ, 1963, стр. 74—79. Ж.

Розенблат-Рот М.

1. Энтропия стохастических процессов, Докл. АН СССР (1957), 112, 1, 16—19. Б, В.
2. Теория передачи информации через стохастические каналы связи, Докл. АН СССР (1957), 112, 2, 202—205. В.
3. Нормированная ϵ -энтропия множеств и передача информации непрерывных источников через непрерывные каналы связи, Докл. АН СССР (1960), 130, 2, 265—268. В.

Рохлин В. А.

1. Об энтропии метрического автоморфизма, Докл. АН СССР (1959), 124, 5, 980—983. Б.

Рубинштейн Г. Ш., Урбаник К.

1. Решение одной экстремальной задачи, Теория вероятн. и ее примен. (1957), 2, 3, 375—377. Г.

Сакагучи М.

1. Заметки по статистическим приложениям теории информации, III, сб. Математика (1959), 3 : 3, 105—115. К.

Сардинас А. А., Патерсон Дж.

1. Необходимое и достаточное условие однозначного разложения закодированных сообщений, Кибернетический сб., вып. 3, М., ИЛ, 1961, стр. 93—102. И.

Сильвермэн Р. А., Болсер М.

1. Кодирование при передаче данных с постоянной скоростью. I. Новый корректирующий код, сб. Коды с обнаружением и исправлением ошибок, М., ИЛ, 1956, стр. 23—43. Ж.

Синай Я. Т.

1. О понятии энтропии динамической системы, Докл. АН СССР (1959), 124, 4, 768—772. Б.
2. О потоках с конечной энтропией, Докл. АН СССР (1959), 125, 6, 1200—1202. Б.
3. Наименьшая ошибка и наилучший способ передачи стационарных сообщений при линейном кодировании и декодировании в случае гауссовских каналов в связи, сб. Проблемы передачи информации, АН СССР (1959), 2, 40—48. Г.

Сифоров В. И.

1. К теории идеального кодирования бинарной передачи, Радиотехн. и электроника (1956), 1, 4, 407—417. Ж.

2. О помехоустойчивости систем с корректирующими кодами, *Радиотехн. и электроника* (1956), 1, 2, 131—142. Е.
3. Параметры систем бинарного кодирования, *Электросвязь* (1957), 1, 3—11. Ж.
4. О наивыгоднейшем использовании кодирующих систем, *Электросвязь* (1957), 5, 7—15. Ж.
5. О пропускной способности каналов связи с медленными случайными изменениями параметров, *Научн. докл. высш. школы, Радиотехн. и электроника* (1958), 1, 3—7. Г.
6. О собственной пропускной способности каналов связи со случайными изменениями параметров, *Научн. докл. высш. школы, Радиотехн. и электроника* (1958), 1, 7—11. Г.
7. Об условиях получения высокой пропускной способности каналов связи со случайными изменениями параметров, *Электросвязь* (1958), 1, 3—8. Г.
8. О пропускной способности каналов связи со случайными изменениями поглощения, *Радиотехн.* (1958), 13, 5, 7—18. Г.

Слепьян Д.

1. Класс двоичных сигнальных алфавитов, сб. Теория передачи сообщений, М., ИЛ, 1957, стр. 82—114. Ж.

Тарасенко Ф. П.

1. Передача информации по марковской цепи, *Тр. Сибирск. физ.-техн. ин-та при Томском ун-те* (1961), 40, 15—17. Б.
2. Об энтропийных характеристиках случайных процессов с непрерывным временем, там же (1961), 40, 24—28. Б.

Тылкин М. Е.

1. О геометрии Хемминга единичных кубов, *Докл. АН СССР* (1960), 134, 5, 1037—1040. Ж.
2. О реализуемости матриц расстояний в единичных кубах, Проблемы кибернетики, вып. 7, М., Физматгиз, 1962, стр. 31—42. Ж.

Фадеев Д. К.

1. К понятию энтропии конечной вероятностной схемы, *Успехи матем. наук* (1956), 11, 1, 227—231. Б.

Файнштейн А.

1. Основы теории информации, М., ИЛ, 1960. А.

Финк Л. М.

1. Пропускная способность симметричных каналов с переменными параметрами при неограниченной полосе частот, *Радиотехн.* (1960), 15, 7, 21—28. Г.
2. О применимости бинарных корректирующих кодов в каналах передачи дискретной информации, *Радиотехн.* (1961), 16, 10, 3—9. Ж.

Финк Л. М., Ковалев Н. И.

1. Распределение вероятностей и энтропийная мощность узкополосного шума с ограниченной амплитудой, *Радиотехн. и электроника* (1960), 5, 7, 1177—1179. Е.

Флейшман Б. С.

1. Конструкция оптимального кода в простейшем случае бинарного канала, *Научн. докл. высш. школы. Радиотехн. и электроника* (1958), 1, 16—21. В, Е.
2. Сравнение трех оптимальных бинарных кодов, имеющих различные способы построения, *Научн. докл. высш. школы, Радиотехн. и электроника* (1958), 1, 58—62. В, Е.
3. О конструктивном доказательстве основной теоремы Шеннона в простейшем бинарном случае (тезисы), Тр. Всес. совещ. по теории вероятн. и матем. статистике, 1958, Ереван, 1960, 66—71. В, Е.
4. Построение оптимального в смысле Шеннона кода в простейшем случае бинарного канала с шумами, сб. тр. Научно-техн. о-ва радио-техн. и электросвязи им. А. С. Попова (1959), 3, 59—95. В, Е.
5. О построении передачи параметра с оптимальным способом декодирования при наличии неаддитивных шумов, *Радиотехн.* (1960), 15, 8, 25—32. В, Е.

Флейшман Б. С., Линковский Г. Б.

1. Максимум энтропии неизвестного дискретного распределения при задании первого момента, *Радиотехн. и электроника* (1958), 3, 4, 554—556. Г.

Харкевич А. А.

1. О наилучшем коде, *Электросвязь* (1956), 2, 65—70. Ж.
2. О теоретически оптимальной системе связи, *Электросвязь* (1957), 5, 15—19. Ж.
3. О ценности информации, Проблемы кибернетики, вып. 4, М., Физматгиз, 1960, стр. 53—59. Б.
4. Асимптотические выражения скорости передачи при высокой надежности, *Радиотехн.* (1962), 17, 1, 76—77. Ж.
5. Очерки общей теории связи, М., Гостехиздат, 1955. А.
6. Одна теорема, относящаяся к корректирующим кодам, *Радиотехн.* (1962), 17, 5, 80. Ж.
7. Простой вывод нижней оценки числа проверочных символов в систематических корректирующих кодах, *Радиотехн.* (1962), 17, 7, 78—79. Ж.

Харкевич А. А., Блох Э. Л.

1. О предельной пропускной способности системы связи, *Радиотехн.* (1955), 10, 2, 74—75. Г.

Хаффмен Д. А.

1. Синтез линейных многотактных кодирующих схем, сб. Теория передачи сообщений, М., ИЛ, 1957, стр. 52—82. Ж.
2. Метод построения кодов с минимальной избыточностью, Кибернетический сб., вып. 3, М., ИЛ, 1961, стр. 79—87. И.

Хинчин А. Я.

1. Понятие энтропии в теории вероятностей, *Успехи матем. наук* (1953), 8, 3, 3—20. Б, В.

2. Об основных теоремах теории информации, *Успехи матем. наук* (1956), 11, 1, 17—75. В.

Х у Г о д и н

1. Три обратные теоремы к теореме Шеннона в теории информации, *Acta math. sinica* (1961), 11, 3, 260—294 (на кит. яз.). В.
2. Об информационной устойчивости последовательности каналов, *Теор. вероятн. и ее примен.* (1962), 7, 3, 271—282. Б.

Х у р г и н Я. И.

1. Оценка пропускной способности некоторых каналов связи со случайно изменяющимися параметрами, *Радиотехн.* (1959), 14, 12, 19—27. Г.

Х э м м и н г Р. В.

1. Коды с обнаружением и исправлением ошибок, сб. Коды с обнаружением и исправлением ошибок, М., ИЛ, 1956, стр. 7—23. Ж.

Ц а р е г р а д с к и й И. П.

1. Замечание о пропускной способности стационарного канала с конечной памятью, *Теория вероятн. и ее примен.* (1958), 3, 1, 84—96; *Успехи матем. наук* (1958), 13, 6, 49—61. В.

Ц з я н Ц з э - п е й

1. Замечание об определении количества информации, *Теория вероятн. и ее примен.* (1958), 3, 1, 99—103. Б.

Ц и р л е р Н.

1. Линейные возвратные последовательности, Кибернетический сб., вып. 6, М., ИЛ, 1963, стр. 55—79. Ж.

Ц ы б а к о в Б. С.

1. О пропускной способности однолучевого канала со случайными изменениями поглощения, *Радиотехн. и электроника* (1959), 42, 44—51. Г.
2. О пропускной способности двухлучевых каналов связи, *Радиотехн. и электроника* (1959), 4, 7, 1116—1123. Г.
3. О пропускной способности каналов с большим числом лучей, *Радиотехн. и электроника* (1959), 4, 9, 1427—1433. Г.
4. Пропускная способность некоторых многолучевых каналов связи, *Радиотехн. и электроника* (1959), 4, 10, 1602—1608. Г.
5. Шенноновская схема для гауссовского сообщения с равномерным спектром и канала с флуктуационным шумом, *Радиотехн. и электроника* (1961), 6, 4, 649—651. В.
6. Линейное кодирование сообщений, *Радиотехн. и электроника* (1962), 7, 1, 25—38. Г.
7. Линейное кодирование изображений, *Радиотехн. и электроника* (1962), 7, 3, 375—385. Г.

Ш а п и р о Г. С., З л о т н и к Д. Л.

1. К математической теории кодов с исправлением ошибок, Кибернетический сб., вып. 5, М., ИЛ, 1962, стр. 7—32. Ж.

Шаров А. И.

1. Идеальный прием сигналов оптимального кода, *Радиотехн. и электроника* (1961), 6, 10, 1595—1600. Ж.

Шастова Г. А.

1. О помехоустойчивости кода Хемминга, *Радиотехн. и электроника* (1958), 3, 1, 19—26. Е, Ж.

Шеннон К. Э.

1. Математическая теория связи, стр. 243 данного сборника. А, Б, В, Г, Ж, И.
2. Теория связи в секретных системах, стр. 333 данного сборника. Л.
3. Современные достижения теории связи, стр. 403 данного сборника. А.
4. Принципы кодово-импульсной модуляции, стр. 414 дан. сб. Л.
5. Связь при наличии шума, стр. 433 данного сборника. А, Г.
6. Некоторые задачи теории информации, стр. 461 данного сборника. А, Б.
7. Пропускная способность канала с шумом при нулевой ошибке, стр. 464 данного сборника. Д, Е.
8. Геометрический подход к теории пропускной способности каналов связи, стр. 488 данного сборника. Г.
9. Каналы с дополнительной информацией на передатчике, стр. 497 данного сборника. Д.
10. Некоторые результаты теории кодирования для каналов с шумами, стр. 509 данного сборника. Г, Е.
11. Замечания о частичном упорядочении каналов связи, стр. 532 данного сборника. Г, Е.
12. Вероятность ошибки для оптимальных кодов в гауссовском канале, стр. 540 данного сборника. Е.
13. Теоремы кодирования для дискретного источника при заданном критерии точности, стр. 587 данного сборника. В.
14. Двусторонние каналы связи, стр. 622 данного сборника. Д.
15. Бандвагон, стр. 667 данного сборника. А.
16. Предсказание и энтропия печатного английского текста, стр. 669 данного сборника. Л.

Шаров А. И.

1. Идеальный прием сигналов оптимального кода, *Радиотехн. и электроника* (1961), 6, 10, 1595—1600. Ж.

Шастова Г. А.

1. О помехоустойчивости кода Хемминга, *Радиотехн. и электроника* (1958), 3, 1, 19—26, Е, Ж.

Шютценбергер Е. Э.

1. О некоторых мерах «информации», используемых в статистике, сб. Теория передачи сообщений, М., ИЛ, 1957; стр. 7—19. Б, К.

Элайес П.

1. Безошибочное кодирование, сб. Коды с обнаружением и исправлением ошибок, М., ИЛ, 1956, стр. 59—71. Е.

2. Кодирование для двух каналов с шумами, сб. Теория передачи сообщений, М., ИЛ, 1957, 114—139. Е.
3. Кодирование в реальных системах связи, Кибернетический сб., вып. 4, М., ИЛ, 1962, стр. 7—30. Л.

Юшкевич А. А.

1. О предельных теоремах, связанных с понятием энтропии цепей Маркова, *Успехи матем. наук* (1953), 8, 5 (57), 177—180. В.

Яглом А. М.

1. Явные формулы для экстраполяции, фильтрации и вычисления количества информации в теории гауссовских стохастических процессов, *Trans. of the second Prague conf. in inform. theory, stat. dec. funct., random. proc., Liblice, 1959, Prague, 1960*, 251—262. Г.

Яглом А. М., Яглом И. М.

1. Вероятность и информация, изд. 2-е, перераб. и доп., Физматгиз, 1960. А.

II. Работы на иностранных языках

Abramson N. M.

1. A class of systematic codes for non-independent errors, *IRE Trans. Inform. Theory* (1959), 5, 4, 150—157. Ж.
2. A note on single error correcting binary codes, *IRE Trans. Inform. Theory* (1960), 6, 502—503. Ж.
3. A partial ordering for binary channels, *IRE Trans. Inform. Theory* (1960), 6, 529—539. Г, Е.

Adler R. L.

1. Ergodic and mixing properties of infinite memory channels, *Proc. Amer. Math. Soc.* (1961), 12, 6, 924—930. В.

Balatoní J., Rényi A.

1. Az entrópia fogal márol, *Magyar tud. akad. Mat. kutató int. közl.* (1956), 1, 1—2, 9—40. В.

Bambah R. P., Joshi D. D., Luthar I. S.

1. Some lower bounds on the number of code points in a minimum distance binary code, I, *Inform. and Control* (1961), 4, 4, 313—319. Ж.
2. Some lower bounds on the number of code points in a minimum distance binary code, II, *Inform. and Control* (1961), 4, 4, 320—323. Ж.

Bannerji R. B.

1. A systematic method for the construction of error-correcting group codes, *Nature* (1960), 186, 4725, 627. Ж.
2. A decoding procedure for double — error correcting Bose — Ray-Chaudhuri codes, *Proc. IRE* (1961), 49, 10, 1585. Ж.
3. On constructing group codes, *Inform. and Control* (1961), 4, 1, 1—4. Ж.

Bar Hillel Y., Carnap R.

1. Semantic information, *Brit. J. Philos. Sci.* (1963), 4, 14, 147—157. Л.
2. An examination of information theory, *Philos. Sci.* (1955), 22, 86—105. Л.

Barnard G. A.

1. Statistical calculation of word entropies for four western language, *IRE Trans. Inform. Theory* (1955), 1, 1, 49—53. Л.

Bartee T. C., Schneider D.

1. An electronic decoder for Bose — Chaudhuri — Hocquenghem error-correcting codes, *IRE Trans. Inform. Theory* (1962), 8, 5, 17—24. Ж.

Bell D. A.

1. Information theory and its engineering application, 2nd. ed., Pitman, London, 1956. А.

Bell D. A., Ross S. C.

1. Negative entropy of Welsh words, *Inform. Theory 3rd London Sympos.*, 1955, London, 1956, 149—153. Л.

Bellman R., Kalada R.

1. Dynamic programming and statistical communication theory, *Proc. Nat. Acad. Sci. USA* (1957), 43, 8, 749—751. Д.
2. On the role of dynamic programming in statistical communication theory, *IRE Trans. Inform. Theory* (1957), 3, 3, 197—203. Д.

Berger L.

1. Quantité d'information et systèmes physiques, *Helv. phys. acta* (1958), 31, 2, 159—166. Л.
2. A note on error detection codes for asymmetric channels, *Inform. and Control* (1961), 4, 1, 68—73. Ж.

Bernstein A. J., Kim W. H.

1. Linear codes for single error correction in symmetric and asymmetric computational processes, *IRE Trans. Inform. Theory* (1962), 8, 1, 29—35. Ж.

Beuz F.

1. Masznahmen zur Störverminderung und Erhöhung der Kanalkapazität bei Impulskodemodulation, *Arch. elektr. Übertrag* (1955), 9, 299—306, 381—387. Г.

Billings A. R.

1. The rate of transmission of information in pulse-code modification systems, *Proc. Inst. Electr. Engrs*, B. C (1958), 105, 444—447. Г.

Billingsley P.

1. Hausdorff dimension in probability theory, I, II. I, *Illinois J. Math.* (1960), 4, 2, 187—209. II, *Illinois J. Math.* (1961), 5, 2, 291—298. Б, Б.
2. On the coding theorem for the noiseless channel, *Ann. Math. Statistics* (1961), 32, 2, 594—601. Б.

Birch J.

1. Approximations for the entropy for functions of Markov chains, *Ann. Math. Statistics* (1962), 33, 3, 930—938. Б.

Bishop W., Buchanan B.

1. Message redundancy vs feedback for reducing message uncertainty, *IRE Internat. Convent. Rec.* (1957), 5, 2, 33—39. Д.

Blachman N.

1. Bounds for entropy, *J. Appl. Phys.* (1953), 24, 10, 1340. Б.
2. Minimum-cost encoding of information, *Trans. IRE* (1954), 3, 139—149. И.
3. The effect of statistically dependent interference upon channel capacity, *IRE Trans. Inform. Theory* (1962), 8, 5, 53—57. Г.
4. On the capacity of band-limited channel perturbed by statistically independent interference, *IRE Trans. Inform. Theory* (1962), 8, 1, 48—55. Г.

Blackwell D.

1. Infinite codes for memoryless channels, *Ann. Math. Statistics* (1959), 30, 1242—1244. Б.

Blackwell D., Breiman L., Thomasian A. J.

1. The capacity of a class of channels, *Ann. Math. Statistics* (1959), 30, 4, 1229—1241. Д.
2. Exponential error bounds for finite state channels, *Proc. 4th Berkely Sympos. Math. Statist and Probability*, 1960, 1, 57—63. Е.
3. The capacities of certain channel classes under random coding, *Ann. Math. Statistics* (1960), 31, 558—567. Д.

Blanc-Lapierre A.

1. Considerations sur la theorie de la transmission de l'information et sur son application a certains domaines de la physique, *Ann. Inst. Henri Poincaré* (1953), 13, 4, 245—296. Л.

Bloom F. J., Chang S. S. L., Harris B., Hauptschein A., Morgan K. C.

1. Improvement of binary transmission by null-zone reception, *Proc. IRE* (1957), 45, 963—975. Г, Е.

Blyth C. R.

1. Note on estimating information, *Ann. Math. Statistics* (1959), 30, 71—79. Г.

Rose R. C., Kneber R. R., Jr.

1. A geometry of binary sequences associated with group alphabets in information theory, *Ann. Math. Statistics* (1960), 31, 1, 113—139. Ж.

Bose R. C., Shrikhande S. S.

1. A note on a result in the theory of code construction, *Inform. and Control* (1959), 2, 183—194. Ж.

Breiman L.

1. The individual ergodic theorem of information theory, *Ann. Math. Statistics* (1947), 28, 3, 809—811. Б.
2. On achieving channel capacity in finite memory channels, *Illinois J. Math.* (1960), 4, 2, 246—252. Б.

Brillouin L.

1. Principe de néguentropie pour l'information, Louis de Broglie, physicien et penseur, Paris, Albin Michel, 1953. Л.

2. The negentropy principle of information, *J. Appl. Phys.* (1953), 24, 9, 1152—1163. Л.
3. Information theory and negative entropy, *Conv. Eletron. Telev.*, Milano, 1954, 1097—1132. Л.
4. Science and information, *Bull. Assoc. Anciens Eleves Ec. Polyt. Fr.* (1956), 59, 51—55. Л.
5. Information theory and the divergent sums in physics, *Ann. Phys. USA* (1958), 5, 3, 243—250. Л.
6. Inevitable experimental errors determinism and information theory, *Inform. and Control* (1959), 2, 45—63. Л.

Brown D. T.

1. Error detecting and correcting binary codes for arithmetic operations, *IRE Trans. Electronic Comput.* (1960), 9, 333—337. Ж.

Brown A. B., Meyers S. T.

1. Evaluation of some error correction methods applicable to digital data transmission, *IRE Donvent. Rec.* (1958), 4, 37. Ж.

Calabi L., Haefeli H. G.

1. A class of binary systematic codes correcting errors at random and in bursts, *IRE Trans. Circuit Theory* (1959), 6, 79—94. Ж.

Calingaert P.

1. Two-dimensional parity checking, *J. Assoc. Comput. Machinery* (1961), 8, 2, 186—200. Ж.

Campopiano C. N.

1. Estimates of entropy of a message source, *Proc. IRE* (1958), 46, 9, 1652. Б.
2. Construction of relatively maximal, systematic codes of specified minimum distance from linear recurring sequences of maximal period, *IRE Trans. Inform. Theory* (1960), 6, 523—539. Ж.
3. Bounds on burst-error correcting codes, там же (1962), 8, 3, 257. Ж.

Capon J.

1. A probabilistic model for runlength coding of pictures, *IRE Trans. Inform. Theory* (1959), 5, 4, 157—163. Д.

Carleson L.

1. Two remarks on the basic theorems of information theory, *Math. Scand.* (1958), 6, 2, 175—180. Б.

Chang S. S. L.

1. Theory of information feedback systems, *IRE Trans. Inform. Theory* (1956), 2, 1, 29—40. Д.
3. Shannon's theory and feedback systems, *J. Basic Engng (Trans. ASME ser. D)*, (1960), 82, 1, 46—50. Д.
2. Capacity of a certain symmetrical binary channel with finite memory, *IRE Trans. Inform. Theory* (1958), 4, 152—158. Г.

Chang S., Harris B., Metzner Y.

1. Optimum message transmission in a finite time, *IRE Trans. Inform. Theory* (1962), 8, 5, 215—224. Д.

Cherry E. C.

1. Introduction to communication theory, New York, Wiley, 1954. A.
2. A history of the theory of information, *Methodos* (1956), 8, 29—30, 57—92. A.

Chien R. T.

1. On the characteristics of error-correcting codes, *IRE Trans. Inform. Theory* (1958), 4, 2, 91. Ж.
2. Orthogonal matrices, error-correcting codes and load-sharing matrix switches, *IRE Trans. Electronic Comput.* (1959), 8, 3, 400. Ж.
3. Group-codes for prescribed error patterns, *IRE Internat. Convent. Rec.* (1960), 8, 4, 125—134. Ж.

Chung K. L.

1. A note on the ergodic theorem of information theory, *Ann. Math. Statistics* (1961), 32, 2, 612—614. B.

Cocke J.

1. Lossless symbol coding with nonprimes, *IRE Trans. Inform. Theory* (1959), 5, 1, 33—34. Ж.

Constantinescu L., Condrea S., Nicolau E.

1. Theoria informatiei, Bucuresti, Ed. tehn. (1958), 211, 6, 50. A.

Corr F.

1. Multiple-burst detection, *Proc. IRE* (1961), 49, 1337. Ж.

Costa de Beauregard O.

1. Sur l'équivalence entre information et entropie et sur l'inversibilité en physique, *C. R. Acad. sci. Paris* (1956), 243, 1728—1730. Л.
2. Sur l'équivalence entre information et entropie dans le rapport $1/k \ln 2$, *C. R. Acad. sci. Paris* (1960), 251, 2898—2900. Л.

Cowell W. R.

1. The use of group codes in error detection and message retransmission, *IRE Trans. Inform. Theory* (1961), 7, 3, 168. Ж.

Crich F. H., Griffith J. S., Orgel L. E.

1. Codes without commas, *Proc. Acad. Sci. USA* (1957), 43, 5, 416—421. Ж.

Csiszár I.

1. Some remarks on the dimension and entropy of random variables, *Acta math. Acad. scient. hung.* (1961), 12, 3—4, 399—408. Б.

Cullman G., Denis-Papin M., Kaufmann L.

1. Elements de calcul informationnel, Paris, Albin Michel, 1960. A.

Davis H., Frautman D. L.

1. The removal of the redundancy due to intersymbol dependence, *IRE Convent. Rec.* (1955), 3, 4, 182—185. И.

Dimsdale B., Weinberg G. M.

1. Programmed error correction in Project Mercury, *Communs Assoc. Comput Mach.* (1960), 3, 12, 649—652. Л.

Dutha J.

1. Some graphical approach to coding problems, *R. C. A. Review* (1957), 18, 4, 466—474. Ж.

Dwork B., Heller R.

1. Results of geometric approach to the theory and construction of non-binary, multiple error and failure correcting codes, *IRE Internat. Convent. Rec.* (1959), 4, 123—129. Ж.

Elias P.

1. Coding for noisy channels, *IRE Internat. Convent. Rec.* (1955), 3, 4, 37—46. E.
2. Predicative coding, Part I, II, *IRE Trans. Inform. Theory* (1955), 1, 16—24, 24—33; *Proc. IRE* (1955), 43, 893. E.
3. List decoding for noisy channels, *IRE WESCON Convent. Rec.* (1957), 2, 94—104. E.
4. Computation in the presence of noise, *IBM J. Res. and Developm.* (1958), 2, 4, 346—253. Л.
5. Coding and information theory, *Revs Mod. Phys.* (1959), 31, 1, 221—226. A.
6. Coding and decoding, Lectures on communication system theory, McGraw-Hill, New York, 1961. A.

Elsas B.

1. A note on P -nary adjacent-error-correcting codes, *IRE Trans. Inform. Theory*, (1960), 6, 1, 13—15. Ж.

Elsas B., Short R. A.

1. A note on optimum burst-error-correcting codes, *IRE Trans. Inform. Theory* (1962), 8, 1, 39—42. Ж.

Epstein M. A.

1. Algebraic decoding for a binary erasure channels, *IRE Internat. Convent. Rec.* (1958), 6, 4, 56—59. 3.

Eshleman V. R.

1. On the wavelength dependence on the information capacity of meteor-burst propagation, *Proc. IRE* (1957), 45, 1710—1714. Г.

Fano R. M.

1. The statistical theory of information, *Nuovo cimento* (1959), 13, 2, 353—372. A.
2. Transmission of information, MIT Press and Wiley, New York — London, 1961 (готовится к печати русский перевод, ИЛ). A.
3. Present trends, Lectures on communication system theory, McGraw-Hill, New York, 1961. A.

Feinstein A.

1. A new basic theorem of information theory, *Trans. IRE* (1954), 4, 2—22. B.
2. Error bounds in noisy channels without memory, *IRE Trans. Inform. Theory* (1955), 1, 2, 13—14. E.
3. On the coding theorem and its converse for finite-memory channels,

Nuovo cimento (1959), 13, 2, 560—575; *Inform. and Control* (1959), 2, 1, 25—44. B.

Flanagan J. E.

1. Coding to achieve Markov type redundancy, *J. Math. and Phys.* (1954), 33, 3, 258—268. B.

Fleischer I.

1. The central concepts of communication theory for infinite alphabets, *J. Math. and Phys.* (1958), 37, 3, 223—228. B.

Flood J. E.

1. Noise-reducing codes for pulse-code modulation, *Proc. Inst. Electr. Eng. B. C.* (1958), 105, 391—397. Д.

Foata

1. Sur la construction des plans factoriels fractionnes et certains codes correcteurs à l'aide des caractères des groupes des abéliens, *Publ. Inst. Statist. Univ. Paris* (1962), II, 1, 55—66. Ж.

Fontaine A. B., Gallager R. G.

1. Error statistics and coding for binary transmission over telephone circuits, *Proc. IRE* (1961), 49, 4, 1059—1065. Д.

Fontaine A. B., Peterson W. W.

1. On coding for the binary symmetric channel, *Trans. AIEE, Comm. and Elec.* (1958), 77, 1, 638—646. Ж.
2. Group code equivalence and optimum codes, *IRE Trans. Inform. Theory* (1959), 5, 60—70. Ж.

Foy W.

1. On instantaneous entropy, *IRE Tr. IT* (1962), 8,5, 267—274. B.

Freiman C. V.

1. Optimal error detection codes for completely asymmetric binary channels, *Inform. and Control* (1962), 5, 1, 64—71. Ж.

Freitag K.

1. Possibilities of reducing redundancy in telephone transmission channels, *Nachrichtentechn. Z.* (1958), 7, 9, 412—418. E.

Fryer R. G.

1. Note on upper bounds for error detecting and error correcting codes of finite length, *IRE Trans. Inform. Theory* (1960), 6, 502. Ж.
2. Analytical development and implementation of an optimum error-correcting code, *Pennsylvania Technol.* (1960), 13, 3, 101—110. Ж.

Gallager R. G.

1. Low density parity-check codes, *IRE Trans. Inform. Theory* (1962), 8, 1, 21—29. Ж.

Garner H. L.

1. Generalized parity checking, *IRE Trans. Electronic Comput.* (1958), 7, 207—213. Ж.

Gilbert E. N.

1. A comparison of signaling alphabets, *Bell System Techn. J.* (1952), 31, 3, 504—522. Ж.

2. Gray codes and paths on the n -cube, *Bell System Techn. J.* (1958), 37, 815—826. Ж.
3. A problem in binary encoding, *Proc. Sympos. Appl. Math.*, 1960, 10, 291—297, American Math. Soc., Providence, R. I., 1960. Ж.
4. Capacity of a burst-noise channel, *Bell System Techn. J.* (1960), 39, 1253—1265. Г.

Gill A.

1. A theorem concerning compact and cyclic sequences, *IRE Trans. Inform. Theory* (1962), 8, 3, 255. Ж.

Golay M. J. E.

1. Binary coding, *IRE Trans. Inform. Theory* (1954), 4, 23—28. Ж.
2. Notes on hybrid coding, *Proc. IRE* (1955), 43, 5, 625. Д.
3. Note on the penny-weighting problem lossless symbol coding with nonprimes etc., *IRE Trans. Inform. Theory* (1958), 4, 103—109. Ж.

Golomb S.

1. A new derivation of the entropy expressions, *IRE Trans. Inform. Theory* (1961), 7, 3, 166—167. Б.

Good I. J., Doog K. C.

1. A paradox concerning rate of information, *Inform. and Control* (1958), 1, 2, 113—126. Г.
2. A paradox concerning rate of information: corrections and additions, *Inform. and Control* (1959), 2, 195—197. Г.

Green P. E. J.

1. A bibliography of soviet literature in noise, correlation and information theory, *IRE Trans. Inform. Theory* (1956), 2, 91—94. А.
2. Information theory in the USSR, *IRE WESCON Convent. Rec.* (1957), 1, 2, 67—83. А.
3. Feedback communication systems. Lectures on communication system theory, McGraw-Hill, New York, 1961. Д.

Grey L. D.

1. Some bounds for error-correcting codes, *IRE Trans. Inform. Theory* (1962), 8, 3, 200—203. Ж.

Griesmer J. H.

1. A bound for error-correcting codes, *IBM J. Res. and Developm.* (1960), 4, 5, 532—542. Ж.

Hagelbarger D. W.

1. Recurrent codes: easily mechanized, burst-correcting, binary codes, *Bell Systems Techn. J.* (1959), 38, 969—984. Ж.

Hajek J.

1. A property of I -divergences of marginal probability distributions, *Чехосл. матем. ж.* (1958), 8, 3, 460—463. Б, К.

Hancock J. C.

1. An introduction to the principles of communication theory, McGraw-Hill, 1961. А.

Harris B., Hauptshein A., Morgan K., Schwartz L.

1. Binary decision feedback systems for maintaining reliability conditions of varying signal strength, *Proc. Natl. Electronics Conf.* 1957, 13, 126—140. Д.

Harris B., Hauptschein A., Schwartz L.

1. Optimum decision feedback systems, *IRE Internat. Convent. Rec.* (1957), 5, 2, 3—10. Д.
2. Minimum cost decision-feedback systems for detecting signals perturbed by additive Gaussian noise, *Operations Res.* (1957), 5, 680—692. Д.

Harris B., Morgan K. C.

1. Binary symmetric decision feedback systems, *Trans. AIEE (Commun. and Electronics)* (1958), 38, 436—443. Д.

Hartmanis J.

1. The application of some basic inequalities for entropy, *Inform. and Control* (1959), 2, 199—213. Б.
2. Linear multivalued sequential coding networks, *IRE Trans. Circuit Theory* (1959), 6, 69—74. Ж.

Hatori H.

1. A note on the entropy of a continuous distribution, *Kodai Math. Sem. Rep.* (1958), 10, 172—176. Б.

Helstrom C. W.

1. Maximum-weight group codes for the balanced M -ary channel, *IRE Trans. Inform. Theory* (1960), 6, 550—555. Ж.
2. Single error-correcting codes for nonbinary balanced channels, *IRE Trans. Inform. Theory* (1961), 7, 1, 2—7. Ж.

Herdan G.

1. An inequality relation between Jule's characteristic K and Shannon's entropy H , *Z. angew. Math. und Phys.* (1958), 9, 1, 69—73. Б, Л.

Hirschman I.

1. A note on entropy, *Amer. J. Math.* (1957), 79, 152—156. Б.

Holmes J. F.

1. Undetected errors in 5-unit code transmission and their elimination, *Computers and Automat.* (1960), 9, 11, 10—13. Ж.

Honda N.

1. The sequential error-correcting code, *Sci. Repts. res. Insts. Tohoku Univ. ser. B-Electr. Commun.* (1956—1957), 8, 3, 113—124. Ж.

Huang R., Johnson R.

1. Information capacity of time-continuous channels, *IRE Trans. Inform. Theory* (1962), 8, 5, 131—138. Г.

Huffman D.

1. A linear circuit viewpoint on error-correcting codes, *IRE Trans. Inform. Theory* (1956), 2, 1, 20—28. Ж.

Ikeda S.

1. Continuity and characterization of Shannon — Wiener information measure for continuous probability distributions, *Ann. Inst. Statist. Math.* (1959), 11, 2, 131—144. Б.

2. A remark on the convergence of Kullback — Leibler's mean information, *Ann. Inst. Statist. Math.* (1960), 12, 1, 81—88. Б, К.
3. A note on the characterization of Shannon — Wiener's measure of information, *Ann. Inst. Statist. Math.* (1962), 13, 3, 259—266. Б.

In g a r d e n R., U r b a n i k K.

1. Information as a fundamental notion of statistical physics, *Bull. Acad. polon. sci. ser. sci. math. astron. et phys.* (1961), 9, 4, 313—316. Л.

I o s i f e s c u M., T h e o d o r e s c u R.

1. Asupra entropiei lanturilor cu legături complete, *Comun. Acad. RPR* (1961), 11, 7, 821—824. Б.

I t o H.

1. On the theory of continuous information, *Proc. Japan Acad.* (1952), 28, 187—191. Б.
2. Principle of the minimum entropy in information theory, *Proc. Japan Acad.* (1953), 29, 194—197. Б.

J a k o b s K.

1. Die Übertragung diskreter Informationen durch periodische und fast-periodische Känale, *Math. Ann.* (1959), 137, 2, 125—135. Б.
2. Einführung in die Informationstheorie, *Z. angew. Math. and Mech.* (1960), 40, 86—94. А.
3. Über Känale vom Dichtetypus, *Math. Z.* (1962), 78, 2, 151—170. Б, В.
4. Über die Struktur der mittleren Entropie, *Math. Z.* (1962), 78, 1, 33—43. Б, В.

J a y n e s E.

1. Information theory and statistical mechanics, *Phys. Rev.* (1957), 108, 2, 171—190. Л.
2. Note on unique decipherability, *IRE Trans. Inform. Theory* (1959), 5, 3, 98—102. И.

J o s h i D. D.

1. L'information en statistique mathématique et dans la théorie des communications, *Publ. Inst. Statist. Univ. Paris* (1959), 8, 2, 81—159. Б, Ж, К.

K a l l i a n p u r G.

1. On the amount of information contained in $A\sigma$ -field, *Contribut. Probability and Statist.*, Stanford Univ. Press, 1960, 265—273. Б.

K a r p R.

1. Minimum-redundancy coding for the discrete noiseless channel, *IRE Trans. Inform. Theory* (1961), 7, 1, 27—38. И.

K a u t z W.

1. Unit-distance error-checking codes, *IRE Trans. Electronic Comput.* (1958), 7, 179—180. Ж.

K e l l o g P. J., K e l l o g D. J.

1. Entropy of information and the odd ball problem, *J. Appl. Phys.* (1954), 25, 1438—1439. Д.

Kelly I.

1. A new interpretation of information rate, *IRE Trans. Inform. Theory* (1956), 2, 3, 185—189. Д.
2. A new interpretation of information rate, *Bell System. Techn. J.* (1956), 35, 4, 917—926, 986. Д.
3. A class of codes for signaling on a noisy continuous channel, *IRE Trans. Inform. Theory* (1960), 6, 1, 22—24. Е, Ж.

Kesten H.

1. Some remarks on the capacity in the continuous case, *Inform. and Control* (1961), 4, 2—3, 169—184. Д.

Kilmer W.

1. Some results on best recurrent-type binary error-correcting codes, *IRE Internat. Convent. Rec.* (1960), 8, 4, 135—147. Ж.
2. Linear-recurrent binary error-correcting codes for memoryless channels, *IRE Trans. Inform. Theory* (1961), 7, 1, 7—13. Ж.

Kim W.

1. Error-correcting codes for an asymmetric nonbinary channel, *IRE Trans. Inform. Theory* (1959), 5, 4, 188—190. Ж.

Kim W., Freiman C.

1. Single error-correcting codes for asymmetric binary channels, *IRE Trans. Inform. Theory* (1959), 5, 2, 62—66. Ж.
2. Multi-error correcting codes for a binary asymmetric channel, *IRE Trans. Circuit Theory* (1959), 6, 71—78. Ж.

Kinney J.

1. Singular functions associated with Markov chains, *Proc. Amer. Math. Soc.* (1958), 9, 4, 602—608. В.

Konheim A.

1. A generalized independence condition and error correction codes, *Amer. Math. Monthly* (1960), 67, 3, 228—231. Ж.

Korezlioglu H.

1. Quantité d'information mutuelle par unite de temps entre deux processus vectoriels gaussiens, stationnaires et stationnairement corrélés, *C.R.Acad. sci. Paris* (1960), 250, 8, 1436—1438. Г.

Kotz S.

1. Exponential bounds on the probability of error for a discrete memoryless channel, *Ann. Math. Statistics* (1961), 32, 2, 577—582. Е.

Kulback S.

1. Certain inequalities in information theory and the Cramer — Rao inequality, *Ann. Math. Statistics* (1954), 25, 4, 745—751. К.
2. An application of information theory to multivariate analysis, *Ann. Math. Statistics* (1956), 2, 27, 122—146. К.
3. Information theory and statistics, New York, Wiley, London, Chapman and Hall, 1959. К.

Kunisawa K.

1. The mathematical foundation of Shannon's information source and

its application to binary coding. (A statistical treatment of binary coding.) *Repts Statist. Applic. Res. Union Japan Scientist and Eng* (1952), 2, 1, 4—26. Б.

Laemmle A.

1. Efficiency of noise reducing codes, *Trans IRE* (1954), 4, 23—28. Ж.

Lee C.

1. Some properties of nonbinary error correcting codes, *IRE Trans. Inform. Theory* (1958), 4, 77—82. Ж.

Leipnik R.

1. Direction of change with refinement for unweighted and weighted information-entropy functionals, *IRE Trans. Inform. Theory* (1959), 5, 4, 184—186. Б.

Lieber N.

1. A note on the mean square weight codes, *Inform. and Control* (1962), 5, 1, 87—89. Г.

Lindley D.

1. On a measure of the information provided by an experiment, *Ann. Math. Statistics* (1956), 27, 4, 986—1005. К.
2. Binomial sampling schemes and the concept of information, *Biometrika* (1957), 44, 1—2, 179—186. К.

Linfoot E.

1. An informational measure of correlation, *Inform. and Control* (1957), 1, 85—89. Б.

Lipp J.

1. Upper bounds for error detecting and correcting codes, *IRE Trans. Inform. Theory* (1960), 6, 557—559. Ж.

Lomnizky Z., Zaremba S.

1. The asymptotic distributions of estimators of the amount of transmitted information, *Inform. and Control* (1959), 2, 3, 266—284. Г.

MacDonald J.

1. Design methods for maximum minimum-distance error-correcting codes, *IBM J. Res. and Developm.* (1960), 4, 1, 43—57. Ж.

Maguire T., Wright E.

1. The examination of error distribution of error-detection and error-correction procedures, *Trans. IRE* (1961), 9, 2, 101. Ж.

Manfrino R.

1. L'entropia della lingua italiana ed il suo calcolo, *Alta frequenza* (1960), 29, 1, 4—29. Л.

Marcovitz A.

1. Sequential generation and decoding of the p -nary Hamming code, *IRE Trans. Inform. Theory* (1961), 7, 1, 53—54. Ж.

Marcus M.

1. Minimum polarized distance codes, *IBM J. Res. and Developm.* (1961), 5, 3, 241—248. Ж.

Mason M.

1. Binary transmission through noise and fading, *IRE Internat. Convent. Rec.* (1957), 5, 2, 69—82. Г, Е.

Mattson H., Solomon G.

1. A new treatment of Bose — Chaudhuri codes, *J. Soc. Industr. and Appl. Math.* (1961), 9, 4, 654—669. Ж.

McCluskey E.

1. Error-correcting codes — A linear programming approach, *Bell System Techn. J.* (1959), 38, 6, 1485—1512. Ж.

McCarthy J.

1. Measures of the value of information, *Proc. Nat. Acad. Sci. USA* (1956), 42, 654—655. Б.

McMillan B.

1. The basic theorems of information theory, *Ann. Math. Statistics* (1953), 24, 2, 196—219. Б.
2. The mathematics of information theory, *IRE Convent. Rec.* (1955), 4, 48—51. Б.
3. Mathematical aspects of information theory, Current trends in information theory, Pittsburgh University Press, 1954. А.
4. A descriptive introduction to the statistical theory of communication, *Nuovo cimento* (1959), 13, 2, 345—352. А.
5. An elementary approach to the theory of information, *SIAM Rev.* (1961), 3, 3, 211—229. А.

Meggit J.

1. Error correcting codes for correcting bursts of errors, *IBM J Res. and Developm.* (1960), 4, 3, 329—334. Ж.
2. Error correcting codes and their implementation for data transmission systems, *IRE Trans. Inform. Theory* (1961), 7, 4, 234—244. Ж.

Melasc.

1. A cyclic code for double error correction, *IBM J Res. and Developm.* (1960), 4, 3, 364—366. Ж.
2. A new group of codes for correction of dependent errors in data transmission, *IBM J. Res. and Developm.* (1960), 4, 1, 58—65. Ж.

Metzner J., Morgan K.

1. Coded binary decision-feedback communication systems, *IRE Trans. Commun. Systems* (1960), 8, 2, 101. Д.

Metzner J., Schwartz L.

1. An extension of the Kelly betting system to binary decision feedback, *Proc. IRE* (1957), 45, 10, 1414—1415. Д.

Miller G., Friedman E.

1. The reconstruction of mutilated English texts, *Inform. and Control* (1957), 1, 1, 38—55. Л.

Mitchell F., Whitehurst R.

1. Further remarks on the odd ball problem as an example in information theory, *J. Appl. Phys.* (1955), 26, 778—779. И.

Mosch A.

1. On the average uncertainty of a continuous probability distribution, *Appl. Scient. Res. R.* (1955), 4, 6, 469—473. Б.

Mourier E.

1. Étude du choix entre deux lois de probabilité, *C. R. Acad. Sci. Paris* (1946), 223, 712—714. К.

Moy Shu-Teh C.

1. Generalizations of Shannon — McMillan theorem, *Pacif. J. Math.* (1961), 11, 2, 705—714. Б.
2. A note on generalizations of Shannon — McMillan theorem, *Pacif. J. Math.* (1961), 11, 4, 1459—1465. Б.

Muroga S.

1. On the capacity of a discrete channel. 1. Mathematical expression of capacity of a channel which is disturbed by noise in its every one symbol and expressible in one state diagram, *J. Phys. Soc. Japan* (1953), 8, 4, 484—494. Г.
2. Sur la capacité de transmission d'un circuit dans une bande de fréquences continue affectée de bruit, *J. Inst. Electr. Commun. Engrs Japan* (1955), 30, 961—964. Г.
3. On the capacity of a discrete channel, *J. Phys. Soc. Japan* (1956), 1103—1120. Г.
4. On the capacity of a noisy continuous channel, *IRE Trans. Inform. Theory* (1957), 3, 1, 44—51, 82. Г.

Nadler M.

1. A 32-point $n = 12$, $d = 5$ code, *IRE Trans. Inform. Theory* (1962), 8, 1, 58. Ж.

Nedomá J.

1. The capacity of a discrete channel, Trans. of the first Prague conf. in inform. theory, stat. dec. funct., random proc., Prague, 1957, 143—182. Б.
2. On non-ergodic channels, Trans. of the second Prague conf. in inform. theory, stat. dec. funct., random proc., Prague, 1960, 363—395. Б.

Neumann P.

1. A note on cyclic permutation error-correcting codes, *Inform. and Control* (1962), 5, 1, 72—86. Ж.
2. On a class of efficient error-limiting variable-length codes, *IRE Trans. Inform. Theory* (1962), 8, 5, 260—266. И.

Norwood L.

1. Upper bound for error-detecting and error-correcting codes, *IRE Trans. Inform. Theory* (1962), 8, 1, 58. Ж.

Oliver B.

1. Efficient coding, *Bell. System Techn. J.* (1952), 31, 4, 724—750. Л.

Parthasarathy K.

1. On the integral representation of the rate of transmission of a stationary channel, *Illinois J. Math.* (1961), 5, 2, 299—305. Б.

Patterson G.

1. Unit-distance number-representation systems, a generalization of the Gray code, *Proc. IRE* (1957), 45, 7, 1024. Ж.

Pérez A.

1. Notions généralisées d'incertitude, d'entropie et d'information du point de vue de la théorie de martingales, Trans. of the first Prague conf. in inform. theory, stat. dec. funct., random proc., Prague, 1957, 183—208. Б, В.
2. Sur la théorie de l'information dans le cas d'un alphabet abstrait, Trans. of the first Prague conf. in inform. theory, stat. dec. funct., random proc., Prague, 1957, 209—244. Б.
3. Sur la convergence des incertitudes entropies et informations échantillons (sample) vers leurs valeurs vraies, Trans. of the first Prague conf. in inform. theory, stat. dec. funct., random proc., Prague, 1957, 245—252. Б.
4. Sur la théorie de l'information et la discernabilité dans les problèmes de décision statistique, Trans. of the second Prague conf. in inform. theory, stat. dec. funct., random proc., Liblice, 1959, Prague, 1960, 413—497. Б.
5. Bounds for error-correcting codes, *IRE Trans. Inform. Theory* (1962), 8, 1, 60. Ж.

Perry K., Wozencraft J.

1. SECO: a self-regulating error correcting coder — decoder, *IRE Trans. Inform. Theory* (1962), 8, 5, 128—135. 3.

Peterson W.

1. An experimental study of a binary code, *Trans. AIEE Comm. and Elec.*, (1958), 77; 388—392. Ж.
2. On checking an adder, *IBM J. Res. and Developm.* (1958), 2, 166—168. Ж.
3. Binary coding for error control, Proc. National Electronics Conference, 1960, 16, 15—21.
4. Error-correcting codes, MIT Press, Wiley New York, London, 1961 (в печати русский перевод, ИЛ). Ж.

Powers K.

1. A prediction theory approach to information rates, *IRE Convent Rec.* (1956), 4, 4, 132—139. Г.

Prange E.

1. The use of information sets in decoding cyclic codes, *IRE Trans. Inform. Theory* (1962), 8, 5, 5—9. Ж.

Price R.

1. On entropy equivalence in the time- and frequency-domains, *Proc. IRE* (1955), 43, 4, 484—485. Г.

Quastler H.

1. Information theory in psychology: Problems and methods, Glencoe, Free Press, 1956. Л

Rajski C.

1. The Bayes rule and the entropy, Trans. of the first Prague conf. of inform. theory, stat. dec. funct. random proc., Prague, 1957, 354. Б.
2. A metric space of discrete probability distributions, *Inform. and Control* (1961), 4, 4. Б.
3. Entropy and metric spaces, Inform. Theory, London, Butterworth Scient. Publs., 1961, 41—43, Discuss., 44—45. Б.

Ramakrishna B., Subramanian R.

1. Relative efficiency of English and German languages for communication of semantic content, *IRE Trans. Inform. Theory* (1958), 4, 3, 127—129. Л

Ratterson G.

1. Unit-distance number-representation systems, a generalization of the Gray code, *Proc. IRE* (1957), 45, 7, 1024. Ж.

Reich E.

1. On the definition of information, *J. Math. Phys.* (1951), 30, 156—161. Б.

Reiger S.

1. Error probabilities of binary data transmission in the presence of random noise, *IRE Convent. Rec.* (1953), 9, 72—79. Е.
2. Error rates in data transmission, *Proc. IRE* (1958), 46, 919—920. Ж.
3. Codes for the correction of «clustered» errors, *IRE Trans. Inform. Theory* (1960), 6, 1, 16—21. Ж.

Rényi A.

1. On a theorem of P. Erdős and its application theory, *Mathematica (RPR)*, (1959), 1, 2, 341—344. Б.
2. On the dimension and entropy of probability distributions, *Acta math. Acad. scient. hung.* (1959), 10, 1—2, 193—215. Б.
3. On measures of entropy and information, Proc. 4th Berkely Sympos. Math. Statist. and Probability, 1960, 1, 547—561. Б.

Rényi A., Balatoni J.

1. Über den Begriff der Entropie, *Math.—Forschungsber.* (1957), 4, 117—134. Б.

Reza F. M.

1. An introduction to information theory, New York — London, McGraw-Hill, 1961. А.

Riekeman E., Glovazky A., McClusky E.

1. Determination of redundancies in a set of patterns, *IRE Trans. Inform. Theory* (1957), 3, 167—168. Л.

Root W.

1. Communications through unspecified additive noise, *Inform. and Control* (1961), 4, 1, 15—29. Д.

Sachs G.

1. Multiple error correction by means of parity checks, *IRE Trans. Inform. Theory* (1958), 4, 4, 145—147. Ж.

Sakaguchi M.

1. Notes on statistical applications of the information theory, *Repts Statist. Applic. Res., Union Japan Scientists and Engrs* (1952), 1, 4, 27—31. К, И.
2. Notes on statistical applications of the information theory, IV, *Repts Statist. Applic. Res., Union Japan Scientists and Engrs* (1959), 6, 54—57. К.
3. Some remarks on the capacity of a communication channel, *J. Operat. Res. Soc. Japan* (1961), 3, 3, 124—132. Г.

Satō H.

1. On the statistics of a code channel, *Rept. Univ. Electro-Commun.* (1954), 6, 21—25. Д.

Schrer E. H.

1. An error-correction code for quaternary data transmission, *Commun. and Electron.* (1961), 55, 231—236. Ж.

Schützenberger M.

1. On an application of semi-group methods to some problems in coding, *IRE Trans. Inform. Theory* (1956), 2, 46—60. Ж.
2. Une théorie algébrique du codage, *C. R. Acad. sci.* (1956), 242, 7, 862—864. Ж.

Schützenberger M., Marcus R.

1. Full decodable code-word sets, *IRE Trans. Inform. Theory* (1959), 5, 1, 12—15. И.

Schwartz L.

1. Capacity of continuous information channel corrected for time or phase jitter, *Proc. IRE* (1961), 49, 2, 513—514. Г.

Segal I.

1. A note on the concept of entropy, *J. Math. and Mech.* (1960), 9, 4, 623—629. Б.

Seidler J.

1. Some remarks on statistical methods in communication theory, *Bull. Acad. Sci. Pol. Sci.* (1957), 5, 4, 261—267. К.
2. Relationships between information theory and decision functions theory, *Trans. of the second Prague conf. in inform. theory, stat. dec. funct., random proc., Liblice, 1959, Prague, 1960*, 579—592. К.

Sellers F.

1. Bit loss and gain correction code, *IRE Trans. Inform. Theory* (1962), 8, 1, 35—39. Ж.

Selmer M.

1. A new upper bound for error-correcting codes, *IRE Trans. Inform. Theory* (1962), 8, 3, 203—208. Ж.

Siforoff W. I.

1. On noise stability of a system with error-correcting codes, *IRE Trans. Inform. Theory* (1956), 24, 109—115, 156. Д.

Silverman R.

1. On binary channels and their cascades, *IRE Trans. Inform. Theory* (1955), 1, 19—27. Г.

Slepian D.

1. A note on two-binary signaling alphabets, *IRE Trans. Inform. Theory* (1956), 2, 2, 84—86, 98. Ж.
2. Coding theory, *Nuovo cimento* (1959), 13, 2, 373—388. Ж.
3. Some further theory of group codes, *Bell System Techn. J.* (1960), 39, 1219—1252. Ж.

Soest J.

1. Some consequences of the finiteness information, Inform. Theory 3rd London Sympos., 1955, 3—7. Б.

Solomon G.

1. A weight formula for group codes, *IRE Trans. Inform. Theory* (1962), 3, 5, 1—4. Ж.

Stahl A.

1. Zur Anwendung des Informationsbegriffes in der statistischen Physik, *Z. Naturforsch.* (1960), 15a, 8, 655—662. Л.

Stam A.

1. On measures of information, *Proc. Nederl. Akad. Wetensch.* (1957), 60, 201—211. Б.
2. Some inequalities satisfied by the quantities of information of Fisher and Shannon, *Inform. and Control* (1959), 2, 101—112. Б, К.

Stern T.

1. Some quantum effects in information channels, *IRE Trans. Inform. Theory* (1960), 6, 435—440. Ж.

Stern T., Friedland B.

1. Application of modular sequential circuits to single-error-correcting p -nary codes, *IRE Trans. Inform. Theory* (1959), 5, 114—123. Ж.

Stone J.

1. Multiple burst error correction, *Inf. and Cont.* (1961) 4, 4, 324—331. Ж.

Storer J.

1. Optimum finite code groups, *Proc. IRE* (1958), 46, 9, 1649. Ж.

Strait P.

1. A lower bound for error-detecting and error-correcting codes, *IRE Trans. Inform. Theory* (1961), 7, 2, 114—118. Ж.

Stumpers F. L.

1. A bibliography of information theory, Communication theory — cybernetics, *IRE Trans.* (1953), PGIT-2, November.
2. A bibliography of information theory, Communication theory — cybernetics (First supplement), *IRE Trans. Inform. Theory* (1955), № 2—3, 31—47. А.
3. A bibliography of information theory, Communication theory — cybernetics (Second supplement), *IRE Trans. Inform. Theory* (1957), 3, 2, 150—166. А.

4. A bibliography of information theory (Third supplement), *IRE Trans. Inform. Theory* (1960), 6, 1, 25—51. A.
- Stutt C.
1. Information rate in a continuous channel for regular-simplex codes, *IRE Trans. Inform. Theory* (1960), 6, 516—523. Г, E.
- Suzuki K.
1. On «amount of information», *Proc. Japan Acad.* (1956), 32, 10, 726—730. B.
 2. On the écart between two «amounts of information», *Proc. Japan Acad.* (1957), 33, 1, 25—28. B.
- Swerling P.
1. Paradoxes related to the rate of transmission of information, *Inform. and Control* (1960), 3, 4, 351—359. Г.
- Takano K.
1. On the basic theorems of information theory, *Ann. Inst. Statist. Math.* (1958), 9, 2, 53—77. B.
- Thomasian A.
1. The metric structure of codes for the binary symmetric channel, *Proc. 4th Berkely Sympos. Math. Statist. and Probability*, 1960, 1, 669—680. E.
 2. An elementary proof of the AEP of information theory, *Ann. Math. Statistics* (1960), 31, 2, 452—456. B.
 3. Error bounds for continuous channels, *Inform. Theory*, London Butterworth Scient. Publs., 1961, 46—59. E.
- Tverberg H.
1. A new derivation of the information function, *Math. scand.* (1958), 6, 2, 297—298. B.
- Turin G.
1. The asymptotic behavior of ideal m -ary systems, *Proc. IRE* (1959), 47, 93. B, E.
- Ulrich W.
1. Non-binary error correction codes, *Bell System Techn. J.* (1957), 36, 6, 1341—1388. Ж.
- Verich W.
1. Non-binary error-correcting codes, *Bell System Techn. J.* (1957), 1, 36, 53—77. Ж.
- Ville J.
1. Leçons sur quelques aspects nouveaux de la théorie des probabilités, *Ann. Inst. Henri Poincaré* (1954), 14, 2, 60—143. A.
- Vincze I.
1. An interpretation of the I -divergence of information theory, *Trans. of the second Prague conf. in inform. theory stat. dec. funct., random proc.*, Liblice, 1959, Prague, 1960. K.
- Votavová L.
1. Ein Satz von Extremen der Entropie, *Trans. of the first Prague Conf.*

in inform. theory stat. dec. funct. random proc., Liblice, Prague, 1957, 293—296. Г.

Watanabe S.

1. A study of ergodicity and redundancy based on intersymbol correlation of finite range, *Trans. IRE* (1954), 4, 85—92. Г.
2. On the information theory and metric lattices, I, II, *Rept. Univ. Electro-Commun.* (1953), 5, 19—38; *Rept. Univ. Electro-Commun.* (1954), 6, 27—35. Б.
3. Binary coding and isometric transformations in a finite metric boolean algebra, *Rept. Univ. Electro-Commun.* (1955), 7, 17—40. Ж.
4. Information theoretical analysis of multivariate correlation, *IBM J. Res. and Developm.* (1960), 4, 1, 66—82. К.

Wax N.

1. On upper bounds for error detecting and error correcting codes of finite length, *IRE Trans. Inform. Theory* (1959), 5, 4, 164—174. Ж.

Weeg G.

1. Uniqueness of weighted code representations, *IRE Trans. Electronic Comput.* (1960), 9, 4, 487—489. Ж.

Weiss L.

1. On the strong converse of the coding theorem for symmetric channels without memory, *Quart. Appl. Math.* (1960), 18, 3, 209—214. Е.

Wolfovitz J.

1. The coding of messages subject to chance errors, *Illinois J. Math.* (1957), 1, 4, 591—606. Б.
2. A upper bound on the rate of transmission of messages, *Illinois J. Math.* (1958), 2, 1, 137—141. Б.
3. Information theory for mathematicians, *Ann. Math. Statistics* (1958), 29, 2, 351—356. Б.
4. The maximum achievable length of an error correcting code, *Illinois J. Math.* (1958), 2, 3, 454—458. Б.
5. La plus grande longueur réalisable d'un code correcteur d'erreurs, Calcul des probabilités et ses applications (Colloques internat. CNRS, 87), 1959, 139—141. Б.
6. Strong converse of the coding theorem for semicontinuous channels, *Illinois J. Math.* (1959), 3, 4, 477—489. Б.
7. A note on the strong converse of the coding theorem for the general discrete finite-memory channel, *Inform. and Control* (1960), 3, 80—93. Б.
8. On coding theorems for simultaneous channels, *IRE Trans. Circuit Theory* (1960), 7, 4, 513—516. Д.
9. Simultaneous channels, *Arch. Ration. Mech. and Analysis* (1960), 4, 4, 371—386. Д.
10. Contributions to information theory, *Proc. Nat. Acad. Sci. USA* (1960), 46, 557—561. Б, Д.
11. On channels in which the distribution of error is known only to the

receiver or only to the sender, Inform. and decision processes, New York — Toronto — London, McGraw-Hill, 1960. Д.

12. A channel with infinite memory, Proc. 4th Berkely Simpos. Math. Statist. and Probability, 1960, 1, 763—767. В.
13. Coding theorems of information theory, Berlin — Göttingen — Heidelberg, Springer, 1961. А.

Wozencraft J.

1. Sequential decoding for reliable communication, *Nat. IRE Convent. Record* (1957), 5, 11—25. 3.
2. Sequential reception of time variant dispersive transmissions. Lectures on communication system theory, McGraw-Hill, New York, 1961. 3.

Zemanek H.

1. Elementare Informationstheorie, Wien — München, Oldenbourg, 1959. А.

Zetterberg L.

1. Cyclic codes from irreducible polynomials for correction of multiple errors, *IRE Trans. Inform. Theory* (1962), 8, 1, 13—21. Ж.

Zierler N.

1. On decoding linear error-correcting codes, *IRE Trans. Inform. Theory* (1960), 6, 450—459. Ж.

Ziv J.

1. Coding and decoding for time-discrete amplitude-continuous memoryless channels, *IRE Trans. Inform. Theory* (1962), 8, 5, 199—205. 3.

ИМЕННОЙ УКАЗАТЕЛЬ

Аккерман (Ackermann W.) 13
Александров А. Д. 488
Александров П. С. 488

Батлер (Butler S.) 162, 179
Беннет (Bennet W.) 413, 437, 588
Беркли (Berkeley E. C.) 59, 179
Беркс (Burks A.) 233
Беттел (Battel) 193
Бигелоу (Bigelou) 226, 224, 231
Биркгоф (Birkhoff G.) 59, 293
Блекуэлл (Blackwell D.) 650, 685
Бодс (Bode H. W.) 322, 463, 687
Больцман (Boltzmann L.) 261, 403
Брейман (Breiman L.) 658
Брозин (Brosin) 230
Буль (Bool G.) 13, 59
Буш (Bush V.) 709
Бэббидж (Babbage Ch.) 162, 163

Васильев Ю. Л. 102
Ватанабе (Watanabe S.) 750
Велч (Welch B. L.) 550, 551, 572
Вернам (Vernam G. S.) 346
Винер (Wiener N.) 180, 215, 293, 294,
309, 322, 403, 687
Виньерон (Vigneron H.) 215
Виттстон (Whittstone) 347
Вуджер (Woodger J.) 59

Габор (Gabor D.) 403, 437
Галлагер (Gallager R. G.) 641
Гейнс (Gaines H. F.) 333
Гельфонд А. О. 248
Герьерри (Guerrieri J.) 711
Гилберт (Gilbert E. N.) 104, 574
Гилиге (Giliege M.) 333
Гильберт (Hilbert D.) 13, 56, 715
Гобсон (Hobson E. W.) 782
Голдстейн (Goldstine H.) 233
Голей (Goley M. J. F.) 290
Гольден (Holden) 715

Греа (Grea R.) 102
Гупта (Gupta) 51, 53, 55
Гуревич (Hurewicz W.) 293, 445

Давенпорт (Davenport W. B.) 291
Данциг (Danzig G.) 730
Деvid (David H.) 558, 571
Девис (Davis M.) 781
Де Гроот (De Groot A. D.) 186, 196,
209, 215
Де-Леу (De Leeuw K.) 751
Дерр (Derr) 196
Джерард (Gerard) 225, 231
Джонсон (Johnson N.) 550, 551,
572
Дибольд (Diebold J.) 179
Добрушин Р. Л. 281, 319, 330, 543,
587, 598
Дуб (Doob J. L.) 293, 343, 781
Дьюи (Dewey G.) 253, 671

Игонне (Higonnet R.) 102

Кантор (Cantor G.) 445
Кардо (Cardot C.) 102
Кейстер (Keister W.) 219
Кельвин (Kelvin) 709
Кемпелен (Kempelen von W.) 182,
193, 217
Кендалл (Kendall M. G.) 251
Кёниг (König D.) 735
Клини (Kleene S. C.) 165, 179, 782
Колмогоров А. Н. 332, 343, 587,
687
Кондон (Condon) 196
Котельников В. А. 295, 435
Котурат (Couturat L.) 13, 59
Крамер (Cramer H.) 638
Кричевский Р. Е. 58
Крускал (Kruskal W.) 558, 571
Крылов А. Н. 709
Кудрявцев Л. Д. 735
Купман (Koorman) 293

Курант (Courant R.) 56
Кэли (Cayley A.) 46
Кэннон (Cannon W. B.) 175

Ландаль (Landahl H. D.) 179
Литльвуд (Littlewood J. E.) 682
Лупанов О. Б. 102, 104
Льюс (Luce R. D.) 341
Любич Ю. И. 322

Мак-Каллок (McCulloch W.) 148, 163, 179, 231
Мак-Карти (McCarthy J.) 233
Мак-Кинси (McKinsey J. C. C.) 341
Мак-Коллум (McCollum D. M.) 167, 179
Мак-Лейн (MacLane S.) 59
Мак-Магон (Mac Mahon P.) 46, 50
Мак-Миллан (McMillan B.) 322, 452
Мецар (Meszar J.) 179
Мид (Mead) 226
Монтгомери (Montgomerie G.) 60
Моргенштерн (Morgenstern O.) 194, 215, 341
Мур (Moore E. F.) 114, 154, 169, 220
Мурога (Muroga S.) 288, 488
Мурский В. Л. 21
Мюирхед (Muirhead) 682

Найквист (Nyquist H.) 243, 403, 434, 437
Накаси́ма (Nakasima A.) 9, 60
Нейман фон (Naumann von J.) 114, 176, 178, 180, 194, 215, 231, 233—239, 293, 341
Новиков П. С. 13

Оливер (Oliver B.) 322, 414, 686
Оруэлл (Orowell G.) 230
Оттингер (Oettinger A. E.) 173, 179

Петерсен (Petersen J.) 736
Пиз (Pease W.) 180
Пинскер М. С. 587
Пирс (Pierce J. R.) 322, 414
Питтс (Pitts W.) 59, 148, 179, 224
Пиш (Piesch H.) 60
Поваров Г. Н. 102, 103
Поля (Полиа) (Polya G.) 682
Пост (Post E. L.) 782
Пратт (Pratt F.) 253, 366, 382, 670

Райс (Rice S. O.) 542, 573
Райт (Wright E. M.) 50, 195, 215

Райффа (Raiffa H.) 341
Риордан (Riordan J.) 46, 82, 83
Россер (Rosser J.) 13
Рут (Root W. L.) 291

Самуэль А. Л. (Samuel A. L.) 220
Сарымсаков Г. А. 256
Севидж (Savage L. J.) 223, 225, 226, 228, 231
Силверман (Silverman R.) 538
Смит (Smith J. B.) 167, 179, 251
Стрэчи (Strachey C. S.) 170, 171, 180, 220, 309
Сулливан (Sullivan H.) 309, 403

Таллер (Tuller W. G.) 309, 403
Тихомиров В. 332
Тойбер (Teuber) 226, 230
Толмен (Tolman R. C.) 261
Томасян (Thomasian A.) 658
Торрес-и-Квеведо (Torres y Quevedo L.) 182, 191, 218
Тоуни (Tawney) 196
Трахтенброт Б. А. 165
Троттер (Trotter H.) 771
Туркетт (Turquette A.) 13
Тьюринг (Turing A. M.) 165, 166, 180, 236, 238, 740, 741, 782

Уайтхед (Whitehead A.) 13
Уивер (Weaver E.) 243
Уиттекер (Whittaker J.) 437
Уолман (Wallman H.) 445

Фаддеев В. К. 261
Файн (Fine R.) 209, 215
Файнштейн (Feinstein A.) 281, 531, 729
Фалкерсон (Fulkerson D.) 730
Фано (Fano R. M.) 272
Фелдман (Feldman C. B.) 413
Феллер (Feller W.) 256, 257, 781
Фёрстер (Foerster von) 228, 230
Финк (Fink D. G.) 413
Форд (Ford L.) 730
Фостер (Foster R. M.) 46, 736
Франк (Frank) 226, 230
Фреше (Frechet M.) 256
Фробениус (Frobenius G.) 525, 531

Халмош (Halmos P. R.) 781
Хантингтон (Huntington E. V.) 14
Харди (Hardy G. N.) 50, 195, 215, 682
Харкевич А. А. 243
Харкнес (Harkness) 193

- Хартли (Hartley R. V. L.) 243, 244, 403, 413, 434
Хартри (Hartree D.) 709
Хаусхолдер (Householder A. S.) 179
Хегельбергер (Hagelbarger D. W.) 174, 221, 651, 652
Хилл (Hill L. S.) 347
Хинчин А. Я. 261, 270
Холбрук (Holbrook B. D.) 110
Хопф (Hopf E.) 293, 398,
Хопф (Hopf H.) 488
Хоэл (Hoel P. G.) 707
Хэмминг (Hamming R.) 289
Ципф (Zipf C. V.) 672
Чандрасекар (Chandrasekhar S.) 250
Чернев (Chernev) 215
Чернов (Chernoff H.) 517, 531, 660
Черч (Church A.) 781
Шапиро (Shapiro N.) 751
Шестаков В. И. 9
Эйлер (Euler L.) 51
Элайес (Elias P.) 487, 531, 543, 729
Эшби (Ashby W. R.) 175, 176, 179
Яблонский С. В. 13, 105
Яглом А. М. 391
Яглом И. М. 391

- Алгебра Буля 13, 59
 - — интерпретация 60
 - — класс функций 92, 128
 - — логическая интерпретация 13
 - — особые правила 61
 - — функция кворума 128
 - секретных систем 350
- Анализатор контактных схем 154
 - — — особенности 161
- Ансамбль функций 291
 - — с ограниченной длительностью 295
 - — — — полосой частот 295
 - — «скорость числа измерений» 33
 - — стационарный 292
 - — эргодический 293
- Асимптотическая надежность 577
- Асимптотическое распределение 705
- Бит 244
- Вероятность ошибки 467, 514, 548
 - — оптимальная 542, 551
 - — — границы 565, 566, 571, 574, 577
 - — средняя для множества кодов 511
 - — — — ансамбля пар кодов 632
- Выходные последовательности случайного устройства 772
- Вычислительные устройства и автоматы 162
- Групповая инвариантность 93
 - — для двух переменных 98
 - — — трех переменных 98
 - — специальный случай 94
- Декодирование 500
 - по минимуму расстояния и максимуму правдоподобия 541
- Декодирующая система 465
 - — оптимальная 541
- Дискретный преобразователь 268
 - Дискретный преобразователь вырожденный 268
 - Дифференциальный анализатор 709
 - — автоматическое управление работой 726
 - — идеализированный 711
 - — передаточные числа 725
 - Избыточность кода 267
 - английского языка 267
 - Инвариантное устройство (оператор) 294
 - Информация взаимная 630
 - мера 245
 - скорость передачи 277, 305, 595
 - средние взаимные скорости передачи 628
 - Искажение 275
 - полное 589
 - Искажения мера 589, 612, 614, 616,
 - — схем 41
 - Источник информации 245, 433
 - — дискретный 249, 457
 - — — с конечным числом состояний 264
 - — непрерывный 458
 - — произведение 596
 - — скорость создания сообщений 318, 320, 458
 - — смешанный 258
 - — с независимыми буквами 591
 - — эргодический 257, 325
 - Исчисление высказываний 14
- Канал 246, 434
 - гауссовский 540
 - двоичный двусторонний умножающий 651
 - — симметричный 543, 611
 - — стирающий 543
 - двусторонний 622, 627
 - — с памятью 658
 - — — симметричной структурой 647

- Канал дискретный без памяти 465, 509
 — — — шума 247
 — — с шумом 275
 — непрерывный 305
 — произведение 466, 524
 — с дополнительной информацией 497
 — — конечным числом состояний 522, 527
 — — обратной связью 479
 — — памятью 620
 — — T концами 660
 — сумма 564, 524
 — частичное упорядочивание 532
 — чистый 534
 Канала надежность 543
 Квантование 415
 Код блоковый для дискретного канала без памяти 509
 — — — двустороннего канала 627
 — — — канала с дополнительной информацией 499
 — входная скорость 465
 — райсовский 573
 — случайный 470
 — Хегельбергера 651
 Кодирование 604, 630, 631
 Кодовое слово 509
 Криптограмма 335, 389

 Логика символическая 13
 — — система аксиом 14

 Матрица смежности 469
 Машина Тьюринга 165, 740
 — — метод построения 741
 — — моделирование 747
 — — понятие A -вычислимости 166
 — — проблема работы 742
 — — с двумя состояниями 741
 — — — одним состоянием 744
 — — универсальная 166
 — — — свойства 166
 Машины вероятностные 751
 — — классы 752
 — вырабатывающие свои принципы игры 220
 — вычислительной и мозга сравнение 237
 — для игры в полностью проанализированные игры 218
 — — — полный анализ которой неизвестен 219
 — — нахождения пути в лабиринте 172

 Машины играющие 168, 216
 — — группы 168
 — — типы 218
 — логические 162
 — обучаемые 171
 — — гомеостат 175
 — стохастические 765
 Модуляция амплитудная 420
 — — однополосная 442
 — — кодово-импульсная 403, 414, 425, 429
 — фазово-импульсная 403, 452
 — частотная 403, 425

 Ненадежность 365
 — ключа 368, 370
 — и избыточность 384
 — распределение 385
 — свойства 370
 — случайного шифра 374
 — сообщения 368, 371

 «Отклонение» вероятности 519
 Отношение сигнал/шум 423
 Отображение, понижающее смежность 474

 Переключатели селекторные 26
 — шаговые 26
 Помехоустойчивость 424
 Пороговая мощность 419
 Пороговой эффект 445
 Предсказание английского языка 673
 — идеальное n -граммное 677, 681
 — линейное 689, 701, 708
 — минимально-квадратичное 707
 — «обратное» 677
 — по методу Винера — Колмогорова 688
 — при наличии шума 701
 Предсказания средне-квадратичная ошибка 694
 Пропускная способность, геометрический подход 488
 — — дискретного канала без памяти с дополнительной информацией 488, 501
 — — — — — шума 247
 — — — — — с шумом 281
 — — — — — канала с конечным числом состояний 522, 527
 — — — — — шумом при нулевой ошибке 464, 466, 467, 469
 — — — — — памятью 621
 — — — — — непрерывного канала 306

- Пропускная способность непрерывно-го канала при ограничении средней мощности 308
- — — — — пиковой мощности 309
 - — при наличии белого шума 446
 - — произвольном типе шума 455
- Пропускной способности область 628, 654, 659
- Процесс марковский 255, 342
- Рабочее число 33
- Реле идеализированные 116
- нижняя оценка числа контактов 142
 - распределение нагрузок 66, 83
 - с неопределенным временем срабатывания 146
 - типы неисправностей 116
- Сглаживание 687, 703
- линейное 689
- Секретные системы 333
- — идеальные 338, 382
 - — идемпотентные 353
 - — подобные 337, 359
 - — совершенно 361, 363
 - — чистые 336
 - — эндоморфные 352
 - — «алгебраические системы» 353
- Система связи 245, 433
- — геометрическое представление 437, 439, 441, 442
 - — дискретная 246
 - — идеальная 451
 - — непрерывная 246
 - — смешанная 246
 - — декодирования 509, 627
 - рабочих чисел 33
- Сети аппроксимация числа 57
- асимптотическое поведение 54
 - максимальный поток через 729
 - параллельно-последовательные 46
 - поведение 47
 - подсчет по Гупта 52
 - — — Дарбу 51
 - — — Эйлеру 52
 - получение производящей функции 48
 - правило соответствия 49
 - приведенные 731
 - расширенные 734
 - ребра 729
 - сечение 731
- Сети существенно-параллельные 47
- существенно-последовательные 47
 - эквивалентность соединений 47
- Сопротивления способы соединения 12
- схем способы определения 65
 - функция 11, 21, 61
- Сумматор 44
- Схем общая теория 19
- переключаательных анализ 59
 - — синтез 59
 - — теория 59
 - сопротивление 10
- Схема гамакообразная 138, 144, 149
- универсальная 38
- Схемы вероятность замыкания 122
- — размыкания 121
 - верхние границы вероятности ошибок 149
 - двоичное сопротивление 10
 - двойственные 30
 - двухполюсные 10, 729
 - длина 121
 - заданной длины и ширины 130
 - метод каскадов 100
 - — нахождения 31
 - — повышения надежности 117
 - — построения 21
 - мостиковые 46
 - направленного действия 148
 - нахождение двойственной 31
 - не параллельно-последовательные 21
 - — — методы построения 21
 - — — многополюсные 19
 - параллельно-последовательные 17
 - параллельные 11
 - переключаательные 9, 154
 - преобразования 19
 - проектирование двухполюсных 155
 - — на анализаторе 155
 - разделительное дерево 89
 - разделительные 67
 - реализующей функцию синтез 32
 - релейные 9
 - — синтез 47
 - — с постоянным напряжением 19
 - с блокировкой 26
 - селекторные 41
 - символический анализ 9
 - синтез надежных 147
 - — переключаательных 47
 - символический анализ 9
 - сложение сопротивлений 11
 - с постоянным током 30

Схемы типы 158

- — соединений 9
- умножение сопротивлений 11
- управления 9
- — свойства 9
- формирование функций нескольких переменных 145
- функция $h(p)$ 120
- ширина 121
- эквивалентной данной, методы нахождения 9

Теорема де Моргана 14

- для канала без шума, основная 270
- — дискретного канала с шумом, основная 281
- о раскраске ребер графа 735
- — схемах и функциях 26—30
- — способах соединения сопротивлений 12, 13
- — эргодических источниках 324
- отсчетов (Котельникова) 435
- синтеза основная 67

Теория автоматов 231

- — вклад фон Неймана 231
- — — — «мультиплекс» 235
- — проблемы эволюции 237

Точность передачи 315

Точности критерий 316, 317

Фильтр линейный 690

- с минимальной фазовой характеристикой 693, 697

Функции несимметрические 33

- произвольной реализация 35
- частично-симметрические, 100
- $h(p)$ свойства 120
- $\lambda(n)$ оценка 82

Функций ансамбль 291

- метод реализации 45

Функций реализация 715

- — симметрических 35, 45
- Функциональная разделимость 93
- Функциональных соотношений типы 93
- Функция переключательная 58
- порождающая 50
- производящая 46, 50
- — получение 48
- сопротивления схемы 61
- $\lambda(n)$ 66
- — поведение 68
- $\mu(n)$ 66

Шифр Бофора 345, 360

- Вернама 346, 364
- Виженера 345, 346, 360, 558
- дробный 348
- Плэйфер 347
- простой подстановки 344
- случайный 374
- чистый 354, 356
- Цезаря 345, 358, 360, 366

Шум белый 292, 304, 408, 409, 448, 454

- — мощность 456
- Гауссовский, см. шум белый
- — произвольный 454
- квантования 415, 423
- тепловой, см. шум белый

Энтропийная мощность 301, 303, 456

Энтропия 261

- английского языка 670
- ансамбля функций 300, 303
- непрерывного распределения 296
- относительная 267
- условная 263, 508

Эффективный процесс 752

Предисловие	5
-----------------------	---

ТЕОРИЯ УПРАВЛЯЮЩИХ СИСТЕМ

Символический анализ релейных и переключательных схем. <i>Перевод Б. Ю. Пильчак</i>	9
Число двухполюсных параллельно-последовательных сетей. <i>Перевод Е. Ю. Захаровой</i>	46
Синтез двухполюсных переключательных схем. <i>Перевод Н. А. Карпова</i>	59
Требования, предъявляемые к объему памяти телефонного коммутатора. <i>Перевод В. В. Мартынюка</i>	106
Надежные схемы из ненадежных реле. <i>Перевод О. Б. Лупанова</i>	114
Использование машины для проектирования переключательных схем. <i>Перевод В. А. Пурто</i>	154
Вычислительные устройства и автоматы. <i>Перевод К. П. Гатченко и М. Г. Гаазе-Рапопорта</i>	162
Машина для игры в шахматы. <i>Перевод И. Б. Задыхайло</i>	181
Составление программ для игры в шахматы на вычислительной машине. <i>Перевод И. Б. Задыхайло</i>	192
Играющие машины. <i>Перевод В. А. Пурто</i>	216
Сообщение о машине, решающей лабиринтную задачу. <i>Перевод М. Г. Белякова и М. Г. Гаазе-Рапопорта</i>	223
Вклад фон Неймана в теорию автоматов. <i>Перевод Е. Ю. Захаровой</i>	232

ТЕОРИЯ ИНФОРМАЦИИ

Математическая теория связи. <i>Перевод В. Ф. Писаренко</i>	243
Теория связи в секретных системах. <i>Перевод В. Ф. Писаренко</i>	333
Современные достижения теории связи. <i>Перевод М. Г. Шура</i>	403
Принципы кодово-импульсной модуляции. <i>Перевод А. А. Харкевича</i>	414
Связь при наличии шума. <i>Перевод А. А. Харкевича</i>	433
Некоторые задачи теории информации. <i>Перевод М. Г. Шура</i>	461
Пропускная способность канала с шумом при нулевой ошибке. <i>Перевод Б. С. Цыбакова</i>	464
Геометрический подход к теории пропускной способности каналов связи. <i>Перевод М. Г. Шура</i>	488

Каналы с дополнительной информацией на передатчике. <i>Перевод Б. С. Цыбакова</i>	497
Некоторые результаты теории кодирования для каналов с шумами. <i>Перевод Б. С. Цыбакова</i>	509
Замечания о частичном упорядочении каналов } связи. <i>Перевод Я. Г. Синая</i>	532
Вероятность ошибки для оптимальных кодов в гауссовском канале. <i>Перевод Б. С. Флейшмана</i>	549
Теоремы кодирования для дискретного источника при заданном критерии точности. <i>Перевод М. С. Пинскера</i>	587
Двусторонние каналы связи. <i>Перевод В. В. Прелова</i>	622

РАЗНОЕ

Бандвагон. <i>Перевод В. Ф. Писаренко</i>	667
Предсказание и энтропия печатного английского текста. <i>Перевод Я. Г. Синая</i>	669
Упрощенный вывод линейной теории сглаживания и предсказания по методу наименьших квадратов. <i>Перевод В. А. Гармаша</i>	687
Математическая теория дифференциального анализатора. <i>Перевод В. В. Мартынюка</i>	709
О максимальном потоке через сеть. <i>Перевод В. И. Левенштейна</i>	729
Теорема о раскраске ребер графа. <i>Перевод О. Б. Лупанова</i>	735
Универсальная машина Тьюринга с двумя внутренними состояниями. <i>Перевод Г. Н. Поварова</i>	740
Вычислимость на вероятностных машинах. <i>Перевод Г. Н. Поварова</i>	751
Библиография	783
Именной указатель	821
Предметный указатель	824

К. ШЕННОН

**РАБОТЫ
ПО ТЕОРИИ ИНФОРМАЦИИ
И
КИБЕРНЕТИКЕ**

Перевод с английского

Под редакцией

Р. Л. ДОБРУШИНА и О. Б. ЛУПАНОВА

С предисловием

А. Н. КОЛМОГОРОВА

ИЗДАТЕЛЬСТВО ИНОСТРАННОЙ ЛИТЕРАТУРЫ

Москва 1963